

Package ‘elastic’

July 22, 2025

Title General Purpose Interface to 'Elasticsearch'

Description Connect to 'Elasticsearch', a 'NoSQL' database built on the 'Java' Virtual Machine. Interacts with the 'Elasticsearch' 'HTTP' API (<<https://www.elastic.co/elasticsearch/>>), including functions for setting connection details to 'Elasticsearch' instances, loading bulk data, searching for documents with both 'HTTP' query variables and 'JSON' based body requests. In addition, 'elastic' provides functions for interacting with API's for 'indices', documents, nodes, clusters, an interface to the cat API, and more.

Version 1.2.0

License MIT + file LICENSE

URL <https://docs.ropensci.org/elastic/> (website),
<https://github.com/ropensci/elastic>

BugReports <https://github.com/ropensci/elastic/issues>

Encoding UTF-8

Language en-US

Imports utils, curl (>= 2.2), crul (>= 0.9.0), jsonlite (>= 1.1), R6

Suggests testthat

RoxygenNote 7.1.1

X-schema.org-applicationCategory Databases

X-schema.org-keywords database, Elasticsearch, HTTP, API, search,
NoSQL, Java, JSON, documents

X-schema.org-isPartOf <https://ropensci.org>

NeedsCompilation no

Author Scott Chamberlain [aut, cre] (ORCID:
<<https://orcid.org/0000-0003-1444-9135>>)

Maintainer Scott Chamberlain <myrmecocystus@gmail.com>

Repository CRAN

Date/Publication 2021-03-16 22:10:05 UTC

Contents

alias	3
cat	5
cluster	10
connect	14
count	16
docs_bulk	18
docs_bulk_create	23
docs_bulk_delete	25
docs_bulk_index	26
docs_bulk_prep	28
docs_bulk_update	32
docs_create	34
docs_delete	36
docs_delete_by_query	37
docs_get	40
docs_mget	42
docs_update	44
docs_update_by_query	46
documents	48
elastic	49
elastic-defunct	51
explain	51
fielddata	53
field_caps	54
field_stats	55
index_template	56
indices	58
ingest	66
mapping	69
msearch	73
mtermvectors	74
nodes	77
percolate	79
ping	85
preference	86
reindex	86
scroll	88
Search	92
searchapis	112
search_shards	113
Search_template	114
Search_uri	117
tasks	122
termvectors	123
tokenizer_set	125
type_remover	127

units-distance	128
units-time	128
validate	129

Index	131
--------------	------------

alias	<i>Elasticsearch alias APIs</i>
-------	---------------------------------

Description

Elasticsearch alias APIs

Usage

```
alias_get(conn, index = NULL, alias = NULL, ignore_unavailable = FALSE, ...)

aliases_get(conn, index = NULL, alias = NULL, ignore_unavailable = FALSE, ...)

alias_exists(conn, index = NULL, alias = NULL, ...)

alias_create(
  conn,
  index,
  alias,
  filter = NULL,
  routing = NULL,
  search_routing = NULL,
  index_routing = NULL,
  ...
)

alias_rename(conn, index, alias, alias_new, ...)

alias_delete(conn, index = NULL, alias, ...)
```

Arguments

conn	an Elasticsearch connection object, see connect()
index	(character) An index name
alias	(character) An alias name
ignore_unavailable	(logical) What to do if an specified index name doesn't exist. If set to TRUE then those indices are ignored.
...	Curl args passed on to curl::verb-POST , curl::verb-GET , curl::verb-HEAD , or curl::verb-DELETE

filter (named list) provides an easy way to create different "views" of the same index. Defined using Query DSL and is applied to all Search, Count, Delete By Query and More Like This operations with this alias. See examples

routing, search_routing, index_routing (character) Associate a routing value with an alias

alias_new (character) A new alias name, used in rename only

Details

Note that you can also create aliases when you create indices by putting the directive in the request body. See the Elasticsearch docs link

Author(s)

Scott Chamberlain myrmecocystus@gmail.com

References

<https://www.elastic.co/guide/en/elasticsearch/reference/current/indices-aliases.html>

Examples

```
## Not run:
# connection setup
(x <- connect())

if (!index_exists(x, "plos")) {
  plosdat <- system.file("examples", "plos_data.json", package = "elastic")
  invisible(docs_bulk(x, plosdat))
}
if (!index_exists(x, "shakespeare")) {
  shake <- system.file("examples", "shakespeare_data.json", package = "elastic")
  invisible(docs_bulk(x, shake))
}

# Create/update an alias
alias_create(x, index = "plos", alias = "candles")
## more than one alias
alias_create(x, index = "plos", alias = c("tables", "chairs"))

# associate an alias with two multiple different indices
alias_create(x, index = c("plos", "shakespeare"), alias = "stools")

# Retrieve a specified alias
alias_get(x, index="plos")
alias_get(x, alias="tables")
alias_get(x, alias="stools")
aliases_get(x)

# rename an alias
```

```

aliases_get(x, "plos")
alias_rename(x, index = 'plos', alias = "stools", alias_new = "plates")
aliases_get(x, "plos")

# filtered aliases
alias_create(x, index = "plos", alias = "candles",
  filter = list(wildcard = list(title = "cell")))
## a search with the alias should give titles with cell in them
(titles <- Search(x, "candles", asdf = TRUE)$hits$hits$`_source.title`)
grepl("cell", titles, ignore.case = TRUE)

# routing
alias_create(x, index = "plos", alias = "candles",
  routing = "1")

# Check for alias existence
alias_exists(x, index = "plos")
alias_exists(x, alias = "tables")
alias_exists(x, alias = "adsfasdf")

# Delete an alias
alias_delete(x, index = "plos", alias = "tables")
alias_exists(x, alias = "tables")

# Curl options
alias_create(x, index = "plos", alias = "tables")
aliases_get(x, alias = "tables", verbose = TRUE)

## End(Not run)

```

cat

Use the cat Elasticsearch api.

Description

Use the cat Elasticsearch api.

Usage

```

cat_(conn, parse = FALSE, ...)

cat_aliases(
  conn,
  verbose = FALSE,
  index = NULL,
  h = NULL,
  help = FALSE,
  bytes = FALSE,
  parse = FALSE,

```

```
    expand_wildcards = "all",  
    ...  
)
```

```
cat_allocation(  
    conn,  
    verbose = FALSE,  
    h = NULL,  
    help = FALSE,  
    bytes = FALSE,  
    parse = FALSE,  
    ...  
)
```

```
cat_count(  
    conn,  
    verbose = FALSE,  
    index = NULL,  
    h = NULL,  
    help = FALSE,  
    bytes = FALSE,  
    parse = FALSE,  
    ...  
)
```

```
cat_segments(  
    conn,  
    verbose = FALSE,  
    index = NULL,  
    h = NULL,  
    help = FALSE,  
    bytes = FALSE,  
    parse = FALSE,  
    ...  
)
```

```
cat_health(  
    conn,  
    verbose = FALSE,  
    h = NULL,  
    help = FALSE,  
    bytes = FALSE,  
    parse = FALSE,  
    ...  
)
```

```
cat_indices(  
    conn,
```

```
        verbose = FALSE,  
        index = NULL,  
        h = NULL,  
        help = FALSE,  
        bytes = FALSE,  
        parse = FALSE,  
        ...  
    )  
  
    cat_master(  
        conn,  
        verbose = FALSE,  
        index = NULL,  
        h = NULL,  
        help = FALSE,  
        bytes = FALSE,  
        parse = FALSE,  
        ...  
    )  
  
    cat_nodes(  
        conn,  
        verbose = FALSE,  
        h = NULL,  
        help = FALSE,  
        bytes = FALSE,  
        parse = FALSE,  
        ...  
    )  
  
    cat_nodeattrs(  
        conn,  
        verbose = FALSE,  
        h = NULL,  
        help = FALSE,  
        bytes = FALSE,  
        parse = FALSE,  
        ...  
    )  
  
    cat_pending_tasks(  
        conn,  
        verbose = FALSE,  
        h = NULL,  
        help = FALSE,  
        bytes = FALSE,  
        parse = FALSE,  
        ...  
    )
```

```
)

cat_plugins(
    conn,
    verbose = FALSE,
    h = NULL,
    help = FALSE,
    bytes = FALSE,
    parse = FALSE,
    ...
)

cat_recovery(
    conn,
    verbose = FALSE,
    index = NULL,
    h = NULL,
    help = FALSE,
    bytes = FALSE,
    parse = FALSE,
    ...
)

cat_thread_pool(
    conn,
    verbose = FALSE,
    index = NULL,
    h = NULL,
    help = FALSE,
    bytes = FALSE,
    parse = FALSE,
    ...
)

cat_shards(
    conn,
    verbose = FALSE,
    index = NULL,
    h = NULL,
    help = FALSE,
    bytes = FALSE,
    parse = FALSE,
    ...
)

cat_fielddata(
    conn,
    verbose = FALSE,
```

```

    index = NULL,
    fields = NULL,
    h = NULL,
    help = FALSE,
    bytes = FALSE,
    parse = FALSE,
    ...
  )

```

Arguments

conn	an Elasticsearch connection object, see connect()
parse	(logical) Parse to a data.frame or not. Default: FALSE
...	Curl args passed on to crul::HttpClient
verbose	(logical) If TRUE (default) the url call used printed to console
index	(character) Index name
h	(character) Fields to return
help	(logical) Output available columns, and their meanings
bytes	(logical) Give numbers back machine friendly. Default: FALSE
expand_wildcards	(character) Whether to expand wildcard expression to concrete indices that are open, closed or both. Valid choices: 'open', 'closed', 'hidden', 'none', 'all'. default: 'all'. Available in ES >= v7.7
fields	(character) Fields to return, only used with fielddata

Details

See <https://www.elastic.co/guide/en/elasticsearch/reference/current/cat.html> for the cat API documentation.

Note how [cat_\(\)](#) has an underscore at the end to avoid conflict with the function [base::cat\(\)](#) in base R.

Examples

```

## Not run:
# connection setup
(x <- connect())

# list Elasticsearch cat endpoints
cat_(x)

# Do other cat operations
cat_aliases(x)
alias_create(x, index = "plos", alias = c("tables", "chairs"))
cat_aliases(x, expand_wildcards='open')
cat_aliases(x, expand_wildcards='all')
cat_allocation(x)

```

```

cat_allocation(x, verbose=TRUE)
cat_count(x)
cat_count(x, index='plos')
cat_count(x, index='gbif')
cat_segments(x)
cat_segments(x, index='gbif')
cat_health(x)
cat_indices(x)
cat_master(x)
cat_nodes(x)
# cat_nodeattrs(x) # not available in older ES versions
cat_pending_tasks(x)
cat_plugins(x)
cat_recovery(x, verbose=TRUE)
cat_recovery(x, index='gbif')
cat_thread_pool(x)
cat_thread_pool(x, verbose=TRUE)
cat_shards(x)
cat_fielddata(x)
cat_fielddata(x, fields='body')

# capture cat data into a data.frame
cat_(x, parse = TRUE)
cat_indices(x, parse = TRUE)
cat_indices(x, parse = TRUE, verbose = TRUE)
cat_count(x, parse = TRUE)
cat_count(x, parse = TRUE, verbose = TRUE)
cat_health(x, parse = TRUE)
cat_health(x, parse = TRUE, verbose = TRUE)

# Get help - what does each column mean
head(cat_indices(x, help = TRUE, parse = TRUE))
cat_health(x, help = TRUE, parse = TRUE)
head(cat_nodes(x, help = TRUE, parse = TRUE))

# Get back only certain fields
cat_nodes(x)
cat_nodes(x, h = c('ip', 'port', 'heapPercent', 'name'))
cat_nodes(x, h = c('id', 'ip', 'port', 'v', 'm'))
cat_indices(x, verbose = TRUE)
cat_indices(x, verbose = TRUE, h = c('index', 'docs.count', 'store.size'))

# Get back machine friendly numbers instead of the normal human friendly
cat_indices(x, verbose = TRUE, bytes = TRUE)

# Curl options
# cat_count(x, timeout_ms = 1)

## End(Not run)

```

Description

Elasticsearch cluster endpoints

Usage

```
cluster_settings(  
    conn,  
    index = NULL,  
    raw = FALSE,  
    callopts = list(),  
    verbose = TRUE,  
    ...  
)  
  
cluster_health(  
    conn,  
    index = NULL,  
    level = NULL,  
    wait_for_status = NULL,  
    wait_for_relocating_shards = NULL,  
    wait_for_active_shards = NULL,  
    wait_for_nodes = NULL,  
    timeout = NULL,  
    raw = FALSE,  
    callopts = list(),  
    verbose = TRUE,  
    ...  
)  
  
cluster_state(  
    conn,  
    index = NULL,  
    metrics = NULL,  
    raw = FALSE,  
    callopts = list(),  
    verbose = TRUE,  
    ...  
)  
  
cluster_stats(  
    conn,  
    index = NULL,  
    raw = FALSE,  
    callopts = list(),  
    verbose = TRUE,  
    ...  
)
```

```
cluster_reroute(conn, body, raw = FALSE, callopts = list(), ...)

cluster_pending_tasks(
  conn,
  index = NULL,
  raw = FALSE,
  callopts = list(),
  verbose = TRUE,
  ...
)
```

Arguments

conn	an Elasticsearch connection object, see connect()
index	Index
raw	If TRUE (default), data is parsed to list. If FALSE, then raw JSON.
callopts	Curl args passed on to crul::verb-POST
verbose	If TRUE (default) the url call used printed to console.
...	Further args passed on to elastic search HTTP API as parameters.
level	Can be one of cluster, indices or shards. Controls the details level of the health information returned. Defaults to cluster.
wait_for_status	One of green, yellow or red. Will wait (until the timeout provided) until the status of the cluster changes to the one provided or better, i.e. green > yellow > red. By default, will not wait for any status.
wait_for_relocating_shards	A number controlling to how many relocating shards to wait for. Usually will be 0 to indicate to wait till all relocations have happened. Defaults to not wait.
wait_for_active_shards	A number controlling to how many active shards to wait for. Defaults to not wait.
wait_for_nodes	The request waits until the specified number N of nodes is available. It also accepts >=N, <=N, >N and <N. Alternatively, it is possible to use ge(N), le(N), gt(N) and lt(N) notation.
timeout	A time based parameter controlling how long to wait if one of the wait_for_XXX are provided. Defaults to 30s.
metrics	One or more of version, master_node, nodes, routing_table, metadata, and blocks. See Details.
body	Query, either a list or json.

Details

metrics param options:

- version Shows the cluster state version.
- master_node Shows the elected master_node part of the response

- `nodes` Shows the nodes part of the response
- `routing_table` Shows the `routing_table` part of the response. If you supply a comma separated list of indices, the returned output will only contain the indices listed.
- `metadata` Shows the metadata part of the response. If you supply a comma separated list of indices, the returned output will only contain the indices listed.
- `blocks` Shows the blocks part of the response

Additional parameters that can be passed in:

- `metric` A comma-separated list of metrics to display. Possible values: `'_all'`, `'completion'`, `'docs'`, `'fielddata'`, `'filter_cache'`, `'flush'`, `'get'`, `'id_cache'`, `'indexing'`, `'merge'`, `'percolate'`, `'refresh'`, `'search'`, `'segments'`, `'store'`, `'warmer'`
- `completion_fields` A comma-separated list of fields for completion metric (supports wildcards)
- `fielddata_fields` A comma-separated list of fields for fielddata metric (supports wildcards)
- `fields` A comma-separated list of fields for fielddata and completion metric (supports wildcards)
- `groups` A comma-separated list of search groups for search statistics
- `allow_no_indices` Whether to ignore if a wildcard indices expression resolves into no concrete indices. (This includes `_all` string or when no indices have been specified)
- `expand_wildcards` Whether to expand wildcard expression to concrete indices that are open, closed or both.
- `ignore_indices` When performed on multiple indices, allows to ignore missing ones (default: `none`)
- `ignore_unavailable` Whether specified concrete indices should be ignored when unavailable (missing or closed)
- `human` Whether to return time and byte values in human-readable format.
- `level` Return stats aggregated at cluster, index or shard level. (`'cluster'`, `'indices'` or `'shards'`, default: `'indices'`)
- `types` A comma-separated list of document types for the indexing index metric

Examples

```
## Not run:
# connection setup
(x <- connect())

cluster_settings(x)
cluster_health(x)

cluster_state(x)
cluster_state(x, metrics = "version")
cluster_state(x, metrics = "nodes")
cluster_state(x, metrics = c("version", "nodes"))
cluster_state(x, metrics = c("version", "nodes", 'blocks'))
cluster_state(x, "shakespeare", metrics = "metadata")
cluster_state(x, c("shakespeare", "flights"), metrics = "metadata")
```

```

cluster_stats(x)
cluster_pending_tasks(x)

body <- '{
  "commands": [
    {
      "move": {
        "index" : "test", "shard" : 0,
        "from_node" : "node1", "to_node" : "node2"
      }
    },
    {
      "allocate_replica" : {
        "index" : "test", "shard" : 1, "node" : "node3"
      }
    }
  ]
}'
# cluster_reroute(x, body = body)

cluster_health(x)
# cluster_health(x, wait_for_status = "yellow", timeout = "3s")

## End(Not run)

```

connect

Set connection details to an Elasticsearch engine.

Description

Set connection details to an Elasticsearch engine.

Usage

```

connect(
  host = "127.0.0.1",
  port = 9200,
  path = NULL,
  transport_schema = "http",
  user = NULL,
  pwd = NULL,
  headers = NULL,
  cainfo = NULL,
  force = FALSE,
  errors = "simple",
  warn = TRUE,
  ignore_version = FALSE,
  ...
)

```

Arguments

host	(character) The base host, defaults to 127.0.0.1
port	(character) port to connect to, defaults to 9200 (optional)
path	(character) context path that is appended to the end of the url. Default: NULL, ignored
transport_schema	(character) http or https. Default: http
user	(character) User name, if required for the connection. You can specify, but ignored for now.
pwd	(character) Password, if required for the connection. You can specify, but ignored for now.
headers	named list of headers. These headers are used in all requests
cainfo	(character) path to a crt bundle, passed to curl option cainfo
force	(logical) Force re-load of connection details. Default: FALSE
errors	(character) One of simple (Default) or complete. Simple gives http code and error message on an error, while complete gives both http code and error message, and stack trace, if available.
warn	(logical) whether to throw warnings from the Elasticsearch server when provided. Pulls warnings from response headers when given. default: TRUE. To turn these off, you can set warn=FALSE or wrap function calls in suppressWarnings() . You can also see warnings in headers by using curl verbose.
ignore_version	(logical) ignore Elasticsearch version checks? default: FALSE. Setting this to TRUE may cause some problems, it has not been fully tested yet. You may want to set this to TRUE if it's not possible to ping the root route of the Elasticsearch instance, which has the Elasticsearch version. We use the version to do alter what request is sent as different Elasticsearch versions allow different parameters.
...	additional curl options to be passed in ALL http requests

Details

The default configuration is set up for localhost access on port 9200, with no username or password. Running this connection method doesn't ping the ES server, but only prints your connection details. All connection details are stored within the returned object. We used to store them in various env vars, but are now contained within the object so you can have any number of connection objects and they shouldn't conflict with one another.

What is the connection object?

Creating a connection object with `connect()` does not create a DBI-like connection object. DBI-like objects have `externalptr`, etc., while `connect()` simply holds details about your Elasticsearch instance (host, port, authentication, etc.) that is used by other methods in this package to interact with your instances' ES API. `connect()` is more or less a fancy list.

You can connect to different Elasticsearch instances within the same R session by creating a separate connection object for each instance; then pass the appropriate connection object to each elastic method.

Examples

```
## Not run:
# the default is set to 127.0.0.1 (i.e., localhost) and port 9200
(x <- connect())
x$make_url()
x$ping()

# pass connection object to function calls
Search(x, q = "*:.*")

# set username/password (hidden in print method)
connect(user = "me", pwd = "stuff")

# set a different host
# connect(host = '162.243.152.53')
# => http://162.243.152.53:9200

# set a different port
# connect(port = 8000)
# => http://localhost:8000

# set a different context path
# connect(path = 'foo_bar')
# => http://localhost:9200/foo_bar

# set to https
# connect(transport_schema = 'https')
# => https://localhost:9200

# set headers
connect(headers = list(a = 'foobar'))

# set cainfo path (hidden in print method)
connect(cainfo = '/some/path/bundle.crt')

## End(Not run)
```

count

Get counts of the number of records per index.

Description

Get counts of the number of records per index.

Usage

```
count(conn, index = NULL, type = NULL, callopts = list(), verbose = TRUE, ...)
```

Arguments

conn	an Elasticsearch connection object, see connect()
index	Index, defaults to all indices
type	Document type, optional
callopts	Curl args passed on to crul::verb-GET
verbose	If TRUE (default) the url call used printed to console.
...	Further args passed on to elastic search HTTP API as parameters.

Details

See docs for the count API here <https://www.elastic.co/guide/en/elasticsearch/reference/current/search-count.html>

You can also get a count of documents using [Search\(\)](#) or [Search_uri\(\)](#) and setting size = 0

Examples

```
## Not run:
# connection setup
(x <- connect())

if (!index_exists(x, "plos")) {
  plosdat <- system.file("examples", "plos_data.json",
    package = "elastic")
  plosdat <- type_remover(plosdat)
  invisible(docs_bulk(x, plosdat))
}
if (!index_exists(x, "shakespeare")) {
  shake <- system.file("examples", "shakespeare_data.json",
    package = "elastic")
  invisible(docs_bulk(x, shake))
}

count(x)
count(x, index='plos')
count(x, index='shakespeare')
count(x, index=c('plos', 'shakespeare'), q="a*")
count(x, index=c('plos', 'shakespeare'), q="z*")

# Curl options
count(x, callopts = list(verbose = TRUE))

## End(Not run)
```

docs_bulk

*Use the bulk API to create, index, update, or delete documents.***Description**

Use the bulk API to create, index, update, or delete documents.

Usage

```
docs_bulk(
  conn,
  x,
  index = NULL,
  type = NULL,
  chunk_size = 1000,
  doc_ids = NULL,
  es_ids = TRUE,
  raw = FALSE,
  quiet = FALSE,
  query = list(),
  digits = NA,
  ...
)
```

Arguments

conn	an Elasticsearch connection object, see connect()
x	A list, data.frame, or character path to a file. required.
index	(character) The index name to use. Required for data.frame input, but optional for file inputs.
type	(character) The type. default: NULL. Note that type is deprecated in Elasticsearch v7 and greater, and removed in Elasticsearch v8
chunk_size	(integer) Size of each chunk. If your data.frame is smaller than chunk_size, this parameter is essentially ignored. We write in chunks because at some point, depending on size of each document, and Elasticsearch setup, writing a very large number of documents in one go becomes slow, so chunking can help. This parameter is ignored if you pass a file name. Default: 1000
doc_ids	An optional vector (character or numeric/integer) of document ids to use. This vector has to equal the size of the documents you are passing in, and will error if not. If you pass a factor we convert to character. Default: not passed
es_ids	(boolean) Let Elasticsearch assign document IDs as UUIDs. These are sequential, so there is order to the IDs they assign. If TRUE, doc_ids is ignored. Default: TRUE
raw	(logical) Get raw JSON back or not. If TRUE you get JSON; if FALSE you get a list. Default: FALSE

quiet	(logical) Suppress progress bar. Default: FALSE
query	(list) a named list of query parameters. optional. options include: pipeline, refresh, routing, _source, _source_excludes, _source_includes, timeout, wait_for_active_shards. See the docs bulk ES page for details
digits	digits used by the parameter of the same name by <code>jsonlite::toJSON()</code> to convert data to JSON before being submitted to your ES instance. default: NA
...	Pass on curl options to <code>curl::HttpClient</code>

Details

More on the Bulk API: <https://www.elastic.co/guide/en/elasticsearch/reference/current/docs-bulk.html>

This function dispatches on data.frame or character input. Character input has to be a file name or the function stops with an error message.

If you pass a data.frame to this function, we by default do an index operation, that is, create the record in the index given by those parameters to the function. Down the road perhaps we will try to support other operations on the bulk API. if you pass a file, of course in that file, you can specify any operations you want.

Row names are dropped from data.frame, and top level names for a list are dropped as well.

A progress bar gives the progress for data.frames and lists - the progress bar is based around a for loop, where progress indicates progress along the iterations of the for loop, where each iteration is a chunk of data that's converted to bulk format, then pushed into Elasticsearch. The character method has no for loop, so no progress bar.

Value

A list

Document IDs

Document IDs can be passed in via the doc_ids paramater when passing in data.frame or list, but not with files. If ids are not passed to doc_ids, we assign document IDs from 1 to length of the object (rows of a data.frame, or length of a list). In the future we may allow the user to select whether they want to assign sequential numeric IDs or to allow Elasticsearch to assign IDs, which are UUIDs that are actually sequential, so you still can determine an order of your documents.

Document IDs and Factors

If you pass in ids that are of class factor, we coerce them to character with `as.character`. This applies to both data.frame and list inputs, but not to file inputs.

Large numbers for document IDs

Until recently, if you had very large integers for document IDs, docs_bulk failed. It should be fixed now. Let us know if not.

Missing data

As of **elastic** version 0.7.8.9515 we convert NA to null before loading into Elasticsearch. Previously, fields that had an NA were dropped - but when you read data back from Elasticsearch into R, you retain those missing values as **jsonlite** fills those in for you. Now, fields with NA's are made into null, and are not dropped in Elasticsearch.

Note also that null values can not be indexed or searched <https://www.elastic.co/guide/en/elasticsearch/reference/5.3/null-value.html>

Tips

This function returns the response from Elasticsearch, but you'll likely not be that interested in the response. If not, wrap your call to docs_bulk in `invisible()`, like so: `invisible(docs_bulk(...))`

Connections/Files

We create temporary files, and connections to those files, when data.frame's and lists are passed in to docs_bulk() (not when a file is passed in since we don't need to create a file). After inserting data into your Elasticsearch instance, we close the connections and delete the temporary files.

There are some exceptions though. When you pass in your own file, whether a tempfile or not, we don't delete those files after using them - in case you need those files again. Your own tempfile's will be cleaned up/delete when the R session ends. Non-tempfile's won't be cleaned up/deleted after the R session ends.

Elasticsearch versions that don't support type

See the `type_remover()` function.

See Also

Other bulk-functions: `docs_bulk_create()`, `docs_bulk_delete()`, `docs_bulk_index()`, `docs_bulk_prep()`, `docs_bulk_update()`

Examples

```
## Not run:
# connection setup
(x <- connect())

# From a file already in newline delimited JSON format
plosdat <- system.file("examples", "plos_data.json", package = "elastic")
docs_bulk(x, plosdat)
aliases_get(x)
index_delete(x, index='plos')
aliases_get(x)

# From a data.frame
docs_bulk(x, mtcars, index = "hello")
## field names cannot contain dots
names(iris) <- gsub("\\.", "_", names(iris))
docs_bulk(x, iris, "iris")
```

```

## type can be missing, but index can not
docs_bulk(x, iris, "flowers")
## big data.frame, 53K rows, load ggplot2 package first
# res <- docs_bulk(x, diamonds, "diam")
# Search(x, "diam")$hits$total

# From a list
docs_bulk(x, apply(iris, 1, as.list), index="iris")
docs_bulk(x, apply(USArrests, 1, as.list), index="arrests")
# dim_list <- apply(diamonds, 1, as.list)
# out <- docs_bulk(x, dim_list, index="diamfromlist")

# When using in a loop
## We internally get last _id counter to know where to start on next bulk
## insert but you need to sleep in between docs_bulk calls, longer the
## bigger the data is
files <- c(system.file("examples", "test1.csv", package = "elastic"),
            system.file("examples", "test2.csv", package = "elastic"),
            system.file("examples", "test3.csv", package = "elastic"))
for (i in seq_along(files)) {
  d <- read.csv(files[[i]])
  docs_bulk(x, d, index = "testes")
  Sys.sleep(1)
}
count(x, "testes")
index_delete(x, "testes")

# You can include your own document id numbers
## Either pass in as an argument
index_create(x, "testes")
files <- c(system.file("examples", "test1.csv", package = "elastic"),
            system.file("examples", "test2.csv", package = "elastic"),
            system.file("examples", "test3.csv", package = "elastic"))
tt <- vapply(files, function(z) NROW(read.csv(z)), numeric(1))
ids <- list(1:tt[1],
            (tt[1] + 1):(tt[1] + tt[2]),
            (tt[1] + tt[2] + 1):sum(tt))
for (i in seq_along(files)) {
  d <- read.csv(files[[i]])
  docs_bulk(x, d, index = "testes", doc_ids = ids[[i]],
            es_ids = FALSE)
}
count(x, "testes")
index_delete(x, "testes")

## or include in the input data
### from data.frame's
index_create(x, "testes")
files <- c(system.file("examples", "test1_id.csv", package = "elastic"),
            system.file("examples", "test2_id.csv", package = "elastic"),
            system.file("examples", "test3_id.csv", package = "elastic"))
readLines(files[[1]])
for (i in seq_along(files)) {

```

```

    d <- read.csv(files[[i]])
    docs_bulk(x, d, index = "testes")
  }
count(x, "testes")
index_delete(x, "testes")

### from lists via file inputs
index_create(x, "testes")
for (i in seq_along(files)) {
  d <- read.csv(files[[i]])
  d <- apply(d, 1, as.list)
  docs_bulk(x, d, index = "testes")
}
count(x, "testes")
index_delete(x, "testes")

# data.frame's with a single column
## this didn't use to work, but now should work
db <- paste0(sample(letters, 10), collapse = "")
index_create(x, db)
res <- data.frame(foo = 1:10)
out <- docs_bulk(x, res, index = db)
count(x, db)
index_delete(x, db)

# data.frame with a mix of actions
## make sure you use a column named 'es_action' or this won't work
## if you need to delete or update you need document IDs
if (index_exists(x, "baz")) index_delete(x, "baz")
df <- data.frame(a = 1:5, b = 6:10, c = letters[1:5], stringsAsFactors = FALSE)
invisible(docs_bulk(x, df, "baz"))
Sys.sleep(3)
(res <- Search(x, 'baz', asdf=TRUE)$hits$hits)
df[1, "a"] <- 99
df[1, "c"] <- "aa"
df[3, "c"] <- 33
df[3, "c"] <- "cc"
df$es_action <- c('update', 'delete', 'update', 'delete', 'delete')
df$id <- res$id
df
invisible(docs_bulk(x, df, "baz", es_ids = FALSE))
### or es_ids = FALSE and pass in document ids to doc_ids
# invisible(docs_bulk(df, "baz", es_ids = FALSE, doc_ids = df$id))
Search(x, 'baz', asdf=TRUE)$hits$hits

# Curl options
plosdat <- system.file("examples", "plos_data.json",
  package = "elastic")
plosdat <- type_remover(plosdat)
invisible(docs_bulk(x, plosdat, verbose = TRUE))

```

```
# suppress progress bar
invisible(docs_bulk(x, mtcars, index = "hello", quiet = TRUE))
## vs.
invisible(docs_bulk(x, mtcars, index = "hello", quiet = FALSE))

## End(Not run)
```

docs_bulk_create

*Use the bulk API to create documents***Description**

Use the bulk API to create documents

Usage

```
docs_bulk_create(
  conn,
  x,
  index = NULL,
  type = NULL,
  chunk_size = 1000,
  doc_ids = NULL,
  es_ids = TRUE,
  raw = FALSE,
  quiet = FALSE,
  query = list(),
  ...
)
```

Arguments

conn	an Elasticsearch connection object, see connect()
x	A list, data.frame, or character path to a file. required.
index	(character) The index name to use. Required for data.frame input, but optional for file inputs.
type	(character) The type. default: NULL. Note that type is deprecated in Elasticsearch v7 and greater, and removed in Elasticsearch v8
chunk_size	(integer) Size of each chunk. If your data.frame is smaller than chunk_size, this parameter is essentially ignored. We write in chunks because at some point, depending on size of each document, and Elasticsearch setup, writing a very large number of documents in one go becomes slow, so chunking can help. This parameter is ignored if you pass a file name. Default: 1000
doc_ids	An optional vector (character or numeric/integer) of document ids to use. This vector has to equal the size of the documents you are passing in, and will error if not. If you pass a factor we convert to character. Default: not passed

es_ids	(boolean) Let Elasticsearch assign document IDs as UUIDs. These are sequential, so there is order to the IDs they assign. If TRUE, doc_ids is ignored. Default: TRUE
raw	(logical) Get raw JSON back or not. If TRUE you get JSON; if FALSE you get a list. Default: FALSE
quiet	(logical) Suppress progress bar. Default: FALSE
query	(list) a named list of query parameters. optional. options include: pipeline, refresh, routing, _source, _source_excludes, _source_includes, timeout, wait_for_active_shards. See the docs bulk ES page for details
...	Pass on curl options to <code>crul::HttpClient</code>

Details

For doing create with a file already prepared for the bulk API, see `docs_bulk()`

Only data.frame's are supported for now.

References

<https://www.elastic.co/guide/en/elasticsearch/reference/current/docs-bulk.html>

See Also

Other bulk-functions: `docs_bulk_delete()`, `docs_bulk_index()`, `docs_bulk_prep()`, `docs_bulk_update()`, `docs_bulk()`

Examples

```
## Not run:
x <- connect()
if (index_exists(x, "foobar")) index_delete(x, "foobar")

df <- data.frame(name = letters[1:3], size = 1:3, id = 100:102)
docs_bulk_create(x, df, 'foobar', es_ids = FALSE)
Search(x, "foobar", asdf = TRUE)$hits$hits

# more examples
docs_bulk_create(x, mtcars, index = "hello")
## field names cannot contain dots
names(iris) <- gsub("\\.", "_", names(iris))
docs_bulk_create(x, iris, "iris")
## type can be missing, but index can not
docs_bulk_create(x, iris, "flowers")
## big data.frame, 53K rows, load ggplot2 package first
# res <- docs_bulk_create(x, diamonds, "diam")
# Search(x, "diam")$hits$total$value

## End(Not run)
```

docs_bulk_delete	<i>Use the bulk API to delete documents</i>
------------------	---

Description

Use the bulk API to delete documents

Usage

```
docs_bulk_delete(
  conn,
  x,
  index = NULL,
  type = NULL,
  chunk_size = 1000,
  doc_ids = NULL,
  raw = FALSE,
  quiet = FALSE,
  query = list(),
  digits = NA,
  ...
)
```

Arguments

conn	an Elasticsearch connection object, see connect()
x	A list, data.frame, or character path to a file. required.
index	(character) The index name to use. Required for data.frame input, but optional for file inputs.
type	(character) The type. default: NULL. Note that type is deprecated in Elasticsearch v7 and greater, and removed in Elasticsearch v8
chunk_size	(integer) Size of each chunk. If your data.frame is smaller than chunk_size, this parameter is essentially ignored. We write in chunks because at some point, depending on size of each document, and Elasticsearch setup, writing a very large number of documents in one go becomes slow, so chunking can help. This parameter is ignored if you pass a file name. Default: 1000
doc_ids	An optional vector (character or numeric/integer) of document ids to use. This vector has to equal the size of the documents you are passing in, and will error if not. If you pass a factor we convert to character. Default: not passed
raw	(logical) Get raw JSON back or not. If TRUE you get JSON; if FALSE you get a list. Default: FALSE
quiet	(logical) Suppress progress bar. Default: FALSE
query	(list) a named list of query parameters. optional. options include: pipeline, refresh, routing, _source, _source_excludes, _source_includes, timeout, wait_for_active_shards. See the docs bulk ES page for details

digits ignored, used in other docs bulk functions but not used here
 ... Pass on curl options to [crul::HttpClient](#)

Details

For doing deletes with a file already prepared for the bulk API, see [docs_bulk\(\)](#)

Only data.frame's are supported for now.

References

<https://www.elastic.co/guide/en/elasticsearch/reference/current/docs-bulk.html>

See Also

Other bulk-functions: [docs_bulk_create\(\)](#), [docs_bulk_index\(\)](#), [docs_bulk_prep\(\)](#), [docs_bulk_update\(\)](#), [docs_bulk\(\)](#)

Examples

```
## Not run:
x <- connect()
if (index_exists(x, "foobar")) index_delete(x, "foobar")

df <- data.frame(name = letters[1:3], size = 1:3, id = 100:102)
invisible(docs_bulk(x, df, 'foobar', es_ids = FALSE))
Search(x, "foobar", asdf = TRUE)$hits$hits

# delete using doc ids from the data.frame you used to create
invisible(docs_bulk_delete(x, df, index = 'foobar'))
Search(x, "foobar", asdf = TRUE)$hits$total$value

# delete by passing in doc ids
## recreate data first
if (index_exists(x, "foobar")) index_delete(x, "foobar")
df <- data.frame(name = letters[1:3], size = 1:3, id = 100:102)
invisible(docs_bulk(x, df, 'foobar', es_ids = FALSE))
docs_bulk_delete(x, df, index = 'foobar', doc_ids = df$id)
Search(x, "foobar", asdf = TRUE)$hits$total$value

## End(Not run)
```

docs_bulk_index

Use the bulk API to index documents

Description

Use the bulk API to index documents

Usage

```
docs_bulk_index(
  conn,
  x,
  index = NULL,
  type = NULL,
  chunk_size = 1000,
  doc_ids = NULL,
  es_ids = TRUE,
  raw = FALSE,
  quiet = FALSE,
  query = list(),
  digits = NA,
  ...
)
```

Arguments

conn	an Elasticsearch connection object, see connect()
x	A list, data.frame, or character path to a file. required.
index	(character) The index name to use. Required for data.frame input, but optional for file inputs.
type	(character) The type. default: NULL. Note that type is deprecated in Elasticsearch v7 and greater, and removed in Elasticsearch v8
chunk_size	(integer) Size of each chunk. If your data.frame is smaller than chunk_size, this parameter is essentially ignored. We write in chunks because at some point, depending on size of each document, and Elasticsearch setup, writing a very large number of documents in one go becomes slow, so chunking can help. This parameter is ignored if you pass a file name. Default: 1000
doc_ids	An optional vector (character or numeric/integer) of document ids to use. This vector has to equal the size of the documents you are passing in, and will error if not. If you pass a factor we convert to character. Default: not passed
es_ids	(boolean) Let Elasticsearch assign document IDs as UUIDs. These are sequential, so there is order to the IDs they assign. If TRUE, doc_ids is ignored. Default: TRUE
raw	(logical) Get raw JSON back or not. If TRUE you get JSON; if FALSE you get a list. Default: FALSE
quiet	(logical) Suppress progress bar. Default: FALSE
query	(list) a named list of query parameters. optional. options include: pipeline, refresh, routing, _source, _source_excludes, _source_includes, timeout, wait_for_active_shards. See the docs bulk ES page for details
digits	digits used by the parameter of the same name by jsonlite::toJSON() to convert data to JSON before being submitted to your ES instance. default: NA
...	Pass on curl options to curl::HttpClient

Details

For doing index with a file already prepared for the bulk API, see `docs_bulk()`

Only data.frame's are supported for now.

References

<https://www.elastic.co/guide/en/elasticsearch/reference/current/docs-bulk.html>

See Also

Other bulk-functions: `docs_bulk_create()`, `docs_bulk_delete()`, `docs_bulk_prep()`, `docs_bulk_update()`, `docs_bulk()`

Examples

```
## Not run:
x <- connect()
if (index_exists(x, "foobar")) index_delete(x, "foobar")

df <- data.frame(name = letters[1:3], size = 1:3, id = 100:102)
docs_bulk_index(x, df, 'foobar')
docs_bulk_index(x, df, 'foobar', es_ids = FALSE)
Search(x, "foobar", asdf = TRUE)$hits$hits

# more examples
docs_bulk_index(x, mtcars, index = "hello")
## field names cannot contain dots
names(iris) <- gsub("\\.", "_", names(iris))
docs_bulk_index(x, iris, "iris")
## type can be missing, but index can not
docs_bulk_index(x, iris, "flowers")
## big data.frame, 53K rows, load ggplot2 package first
# res <- docs_bulk_index(x, diamonds, "diam")
# Search(x, "diam")$hits$total$value

## End(Not run)
```

docs_bulk_prep

Use the bulk API to prepare bulk format data

Description

Use the bulk API to prepare bulk format data

Usage

```
docs_bulk_prep(
  x,
  index,
  path,
  type = NULL,
  chunk_size = 1000,
  doc_ids = NULL,
  quiet = FALSE,
  digits = NA
)
```

Arguments

x	A data.frame or a list. required.
index	(character) The index name. required.
path	(character) Path to the file. If data is broken into chunks, we'll use this path as the prefix, and suffix each file path with a number. required.
type	(character) The type. default: NULL. Note that type is deprecated in Elasticsearch v7 and greater, and removed in Elasticsearch v8
chunk_size	(integer) Size of each chunk. If your data.frame is smaller than chunk_size, this parameter is essentially ignored. We write in chunks because at some point, depending on size of each document, and Elasticsearch setup, writing a very large number of documents in one go becomes slow, so chunking can help. This parameter is ignored if you pass a file name. Default: 1000
doc_ids	An optional vector (character or numeric/integer) of document ids to use. This vector has to equal the size of the documents you are passing in, and will error if not. If you pass a factor we convert to character. Default: not passed
quiet	(logical) Suppress progress bar. Default: FALSE
digits	digits used by the parameter of the same name by jsonlite::toJSON() to convert data to JSON before being submitted to your ES instance. default: NA

Value

File path(s). By default we use temporary files; these are cleaned up at the end of a session

Tempfiles

In docs_bulk we create temporary files in some cases, and delete those before the function exits. However, we don't clean up those files in this function because the point of the function is to create the newline delimited JSON files that you need. Tempfiles are cleaned up when you R session ends though - be aware of that. If you want to keep the files make sure to move them outside of the temp directory.

See Also

Other bulk-functions: [docs_bulk_create\(\)](#), [docs_bulk_delete\(\)](#), [docs_bulk_index\(\)](#), [docs_bulk_update\(\)](#), [docs_bulk\(\)](#)

Examples

```
## Not run:
# From a data.frame
ff <- tempfile(fileext = ".json")
docs_bulk_prep(mtcars, index = "hello", path = ff)
readLines(ff)

## field names cannot contain dots
names(iris) <- gsub("\\.", "_", names(iris))
docs_bulk_prep(iris, "iris", path = tempfile(fileext = ".json"))

## type can be missing, but index can not
docs_bulk_prep(iris, "flowers", path = tempfile(fileext = ".json"))

# From a list
docs_bulk_prep(apply(iris, 1, as.list), index="iris",
  path = tempfile(fileext = ".json"))
docs_bulk_prep(apply(USArrests, 1, as.list), index="arrests",
  path = tempfile(fileext = ".json"))

# when chunking
## multiple files created, one for each chunk
bigiris <- do.call("rbind", replicate(30, iris, FALSE))
docs_bulk_prep(bigiris, index = "big", path = tempfile(fileext = ".json"))

# When using in a loop
## We internally get last _id counter to know where to start on next bulk
## insert but you need to sleep in between docs_bulk_prep calls, longer the
## bigger the data is
files <- c(system.file("examples", "test1.csv", package = "elastic"),
  system.file("examples", "test2.csv", package = "elastic"),
  system.file("examples", "test3.csv", package = "elastic"))
paths <- vector("list", length = length(files))
for (i in seq_along(files)) {
  d <- read.csv(files[[i]])
  paths[i] <- docs_bulk_prep(d, index = "stuff",
    path = tempfile(fileext = ".json"))
}
unlist(paths)

# You can include your own document id numbers
## Either pass in as an argument
files <- c(system.file("examples", "test1.csv", package = "elastic"),
  system.file("examples", "test2.csv", package = "elastic"),
  system.file("examples", "test3.csv", package = "elastic"))
tt <- vapply(files, function(z) NROW(read.csv(z)), numeric(1))
ids <- list(1:tt[1],
  (tt[1] + 1):(tt[1] + tt[2]),
  (tt[1] + tt[2] + 1):sum(tt))
paths <- vector("list", length = length(files))
for (i in seq_along(files)) {
  d <- read.csv(files[[i]])
```

```

    paths[i] <- docs_bulk_prep(d, index = "testes",
                               doc_ids = ids[[i]], path = tempfile(fileext = ".json"))
  }
  unlist(paths)

## or include in the input data
### from data.frame's
files <- c(system.file("examples", "test1_id.csv", package = "elastic"),
            system.file("examples", "test2_id.csv", package = "elastic"),
            system.file("examples", "test3_id.csv", package = "elastic"))
paths <- vector("list", length = length(files))
for (i in seq_along(files)) {
  d <- read.csv(files[[i]])
  paths[i] <- docs_bulk_prep(d, index = "testes",
                             path = tempfile(fileext = ".json"))
}
unlist(paths)

### from lists via file inputs
paths <- vector("list", length = length(files))
for (i in seq_along(files)) {
  d <- read.csv(files[[i]])
  d <- apply(d, 1, as.list)
  paths[i] <- docs_bulk_prep(d, index = "testes",
                             path = tempfile(fileext = ".json"))
}
unlist(paths)

# A mix of actions
## make sure you use a column named 'es_action' or this won't work
## if you need to delete or update you need document IDs
if (index_exists(x, "baz")) index_delete(x, "baz")
df <- data.frame(a = 1:5, b = 6:10, c = letters[1:5], stringsAsFactors = FALSE)
f <- tempfile(fileext = ".json")
invisible(docs_bulk_prep(df, "baz", f))
cat(readLines(f), sep = "\n")
docs_bulk(x, f)
Sys.sleep(2)
(res <- Search(x, 'baz', asdf=TRUE)$hits$hits)

df[1, "a"] <- 99
df[1, "c"] <- "aa"
df[3, "c"] <- 33
df[3, "c"] <- "cc"
df$es_action <- c('update', 'delete', 'update', 'delete', 'delete')
df$id <- res$id
df
f <- tempfile(fileext = ".json")
invisible(docs_bulk_prep(df, "baz", path = f, doc_ids = df$id))
cat(readLines(f), sep = "\n")
docs_bulk(x, f)

```

```
# suppress progress bar
docs_bulk_prep(mtcars, index = "hello",
  path = tempfile(fileext = ".json"), quiet = TRUE)
## vs.
docs_bulk_prep(mtcars, index = "hello",
  path = tempfile(fileext = ".json"), quiet = FALSE)

## End(Not run)
```

docs_bulk_update

Use the bulk API to update documents

Description

Use the bulk API to update documents

Usage

```
docs_bulk_update(
  conn,
  x,
  index = NULL,
  type = NULL,
  chunk_size = 1000,
  doc_ids = NULL,
  raw = FALSE,
  quiet = FALSE,
  query = list(),
  digits = NA,
  ...
)
```

Arguments

conn	an Elasticsearch connection object, see connect()
x	A list, data.frame, or character path to a file. required.
index	(character) The index name to use. Required for data.frame input, but optional for file inputs.
type	(character) The type. default: NULL. Note that type is deprecated in Elasticsearch v7 and greater, and removed in Elasticsearch v8
chunk_size	(integer) Size of each chunk. If your data.frame is smaller than chunk_size, this parameter is essentially ignored. We write in chunks because at some point, depending on size of each document, and Elasticsearch setup, writing a very large number of documents in one go becomes slow, so chunking can help. This parameter is ignored if you pass a file name. Default: 1000

doc_ids	An optional vector (character or numeric/integer) of document ids to use. This vector has to equal the size of the documents you are passing in, and will error if not. If you pass a factor we convert to character. Default: not passed
raw	(logical) Get raw JSON back or not. If TRUE you get JSON; if FALSE you get a list. Default: FALSE
quiet	(logical) Suppress progress bar. Default: FALSE
query	(list) a named list of query parameters. optional. options include: pipeline, refresh, routing, _source, _source_excludes, _source_includes, timeout, wait_for_active_shards. See the docs bulk ES page for details
digits	digits used by the parameter of the same name by <code>jsonlite::toJSON()</code> to convert data to JSON before being submitted to your ES instance. default: NA
...	Pass on curl options to <code>curl::HttpClient</code>

Details

- doc_as_upsert - is set to TRUE for all records

For doing updates with a file already prepared for the bulk API, see `docs_bulk()`

Only data.frame's are supported for now.

References

<https://www.elastic.co/guide/en/elasticsearch/reference/current/docs-bulk.html#bulk-update>

See Also

Other bulk-functions: `docs_bulk_create()`, `docs_bulk_delete()`, `docs_bulk_index()`, `docs_bulk_prep()`, `docs_bulk()`

Examples

```
## Not run:
x <- connect()
if (index_exists(x, "foobar")) index_delete(x, "foobar")

df <- data.frame(name = letters[1:3], size = 1:3, id = 100:102)
invisible(docs_bulk(x, df, 'foobar', es_ids = FALSE))

# add new rows in existing fields
(df2 <- data.frame(size = c(45, 56), id = 100:101))
(df2 <- data.frame(size = c(45, 56)))
df2$`_id` <- 100:101
df2
Search(x, "foobar", asdf = TRUE)$hits$hits
invisible(docs_bulk_update(x, df2, index = 'foobar'))
Search(x, "foobar", asdf = TRUE)$hits$hits

# add new fields (and new rows by extension)
(df3 <- data.frame(color = c("blue", "red", "green"), id = 100:102))
```

```

Search(x, "foobar", asdf = TRUE)$hits$hits
invisible(docs_bulk_update(x, df3, index = 'foobar'))
Sys.sleep(2) # wait for a few sec to make sure you see changes reflected
Search(x, "foobar", asdf = TRUE)$hits$hits

## End(Not run)

```

docs_create

Create a document

Description

Create a document

Usage

```

docs_create(
  conn,
  index,
  body,
  type = NULL,
  id = NULL,
  version = NULL,
  version_type = NULL,
  op_type = NULL,
  routing = NULL,
  parent = NULL,
  timestamp = NULL,
  ttl = NULL,
  refresh = NULL,
  timeout = NULL,
  callopts = list(),
  ...
)

```

Arguments

conn	an Elasticsearch connection object, see connect()
index	(character) The name of the index. Required
body	The document
type	(character) The type of the document. optional
id	(numeric/character) The document ID. Can be numeric or character. Optional. if not provided, Elasticsearch creates the ID for you as a UUID.
version	(character) Explicit version number for concurrency control
version_type	(character) Specific version type. One of internal, external, external_gte, or force

op_type	(character) Operation type. One of create, or ...
routing	(character) Specific routing value
parent	(numeric) A parent document ID
timestamp	(date) Explicit timestamp for the document
ttl	(aka “time to live”) Expiration time for the document. Expired documents will be expunged automatically. The expiration date that will be set for a document with a provided ttl is relative to the timestamp of the document, meaning it can be based on the time of indexing or on any time provided. The provided ttl must be strictly positive and can be a number (in milliseconds) or any valid time value (e.g, 86400000, 1d).
refresh	(logical) Refresh the index after performing the operation
timeout	(character) Explicit operation timeout, e.g., 5m (for 5 minutes)
callopts	Curl options passed on to <code>crul::HttpClient</code>
...	Further args to query DSL

References

https://www.elastic.co/guide/en/elasticsearch/reference/current/docs-index_.html

Examples

```
## Not run:
(x <- connect())

if (!index_exists(x, 'plos')) {
  plosdat <- system.file("examples", "plos_data.json",
    package = "elastic")
  plosdat <- type_remover(plosdat)
  invisible(docs_bulk(x, plosdat))
}

# give a document id
z <- docs_create(x, index = 'plos', id = 1002,
  body = list(id = "12345", title = "New title"))
z
# and the document is there now
docs_get(x, index = 'plos', id = 1002)

# let Elasticsearch create the document id for you
z <- docs_create(x, index='plos', body=list(id="6789", title="Some title"))
z
# and the document is there now
docs_get(x, index='plos', id=z$id)

## End(Not run)
```

docs_delete*Delete a document*

Description

Delete a document

Usage

```
docs_delete(  
    conn,  
    index,  
    id,  
    type = NULL,  
    refresh = NULL,  
    routing = NULL,  
    timeout = NULL,  
    version = NULL,  
    version_type = NULL,  
    callopts = list(),  
    ...  
)
```

Arguments

conn	an Elasticsearch connection object, see connect()
index	(character) The name of the index. Required
id	(numeric/character) The document ID. Can be numeric or character. Required
type	(character) The type of the document. optional
refresh	(logical) Refresh the index after performing the operation
routing	(character) Specific routing value
timeout	(character) Explicit operation timeout, e.g., 5m (for 5 minutes)
version	(character) Explicit version number for concurrency control
version_type	(character) Specific version type. One of internal or external
callopts	Curl args passed on to curl::HttpClient
...	Further args to query DSL

References

<https://www.elastic.co/guide/en/elasticsearch/reference/current/docs-delete.html>

Examples

```
## Not run:
(x <- connect())
x$ping()

if (!index_exists(x, "plos")) {
  plosdat <- system.file("examples", "plos_data.json",
    package = "elastic")
  plosdat <- type_remover(plosdat)
  docs_bulk(x, plosdat)
}

# delete a document
if (!docs_get(x, index='plos', id=36, exists=TRUE)) {
  docs_create(x, index='plos', id=36,
    body = list(id="12345", title="New title")
  )
}
docs_get(x, index='plos', id=36)
docs_delete(x, index='plos', id=36)
# docs_get(x, index='plos', id=36) # and the document is gone

## End(Not run)
```

docs_delete_by_query *Delete documents by query*

Description

delete documents by query via a POST request

Usage

```
docs_delete_by_query(
  conn,
  index,
  body,
  type = NULL,
  conflicts = NULL,
  routing = NULL,
  scroll_size = NULL,
  refresh = NULL,
  wait_for_completion = NULL,
  wait_for_active_shards = NULL,
  timeout = NULL,
  scroll = NULL,
  requests_per_second = NULL,
  ...
)
```

Arguments

conn	an Elasticsearch connection object, see connect()
index	(character) The name of the index. Required
body	(character/json) query to be passed on to POST request body
type	(character) The type of the document. optional
conflicts	(character) If you'd like to count version conflicts rather than cause them to abort then set conflicts=proceed
routing	(character) Specific routing value
scroll_size	(integer) By default uses scroll batches of 1000. Change batch size with this parameter.
refresh	(logical) Refresh the index after performing the operation
wait_for_completion	(logical) If wait_for_completion=FALSE then Elasticsearch will perform some preflight checks, launch the request, and then return a task which can be used with Tasks APIs to cancel or get the status of the task. Elasticsearch will also create a record of this task as a document at .tasks/task/\$taskId. This is yours to keep or remove as you see fit. When you are done with it, delete it so Elasticsearch can reclaim the space it uses. Default: TRUE
wait_for_active_shards	(logical) controls how many copies of a shard must be active before proceeding with the request.
timeout	(character) Explicit operation timeout, e.g., 5m (for 5 minutes)
scroll	(integer) control how long the "search context" is kept alive, eg scroll='10m', by default it's 5 minutes (5m)
requests_per_second	(integer) any positive decimal number (1.4, 6, 1000, etc); throttles rate at which _delete_by_query issues batches of delete operations by padding each batch with a wait time. The throttling can be disabled by setting requests_per_second=-1
...	Curl args passed on to curl::verb-POST

References

<https://www.elastic.co/guide/en/elasticsearch/reference/current/docs-delete-by-query.html>

See Also

[docs_update_by_query\(\)](#)

Examples

```
## Not run:
(x <- connect())
x$ping()
```

```

plosdat <- system.file("examples", "plos_data.json",
  package = "elastic")
plosdat <- type_remover(plosdat)
if (!index_exists(x, "plos")) invisible(docs_bulk(x, plosdat))

# delete with fuzzy matching
body <- '{
  "query": {
    "match": {
      "title": {
        "query": "cells",
        "fuzziness": 1
      }
    }
  }
}'
docs_delete_by_query(x, index='plos', body = body)

# delete with no fuzziness
if (index_exists(x, "plos")) index_delete(x, 'plos')
invisible(docs_bulk(x, plosdat))
count(x, "plos")
body <- '{
  "query": {
    "match": {
      "title": {
        "query": "cells",
        "fuzziness": 0
      }
    }
  }
}'
docs_delete_by_query(x, index='plos', body = body)

# delete all docs with match_all query
if (index_exists(x, "plos")) index_delete(x, 'plos')
invisible(docs_bulk(x, plosdat))
body <- '{
  "query": {
    "match_all": {}
  }
}'
docs_delete_by_query(x, index='plos', body = body)

# put plos back in
if (index_exists(x, "plos")) index_delete(x, 'plos')
invisible(docs_bulk(x, plosdat))

# delete docs from more than one index
foo <- system.file("examples/foo.json", package = "elastic")
if (!index_exists(x, "foo")) invisible(docs_bulk(x, foo))
bar <- system.file("examples/bar.json", package = "elastic")
if (!index_exists(x, "bar")) invisible(docs_bulk(x, bar))

```

```

body <- '{
  "query": {
    "match_all": {}
  }
}'
docs_delete_by_query(x, index=c('foo', 'bar'),
  body = body, verbose = TRUE)

## End(Not run)

```

docs_get

*Get documents***Description**

Get documents

Usage

```

docs_get(
  conn,
  index,
  id,
  type = NULL,
  source = NULL,
  fields = NULL,
  source_includes = NULL,
  source_excludes = NULL,
  exists = FALSE,
  raw = FALSE,
  callopts = list(),
  verbose = TRUE,
  ...
)

```

Arguments

conn	an Elasticsearch connection object, see connect()
index	(character) The name of the index. Required
id	(numeric/character) The document ID. Can be numeric or character. Required
type	(character) The type of the document. optional
source	(logical) If TRUE (default), return source. note that it is actually set to NULL in the function definition, but within Elasticsearch, it returns the source by default. alternatively, you can pass a vector of field names to return.
fields	Fields to return from the response object.

source_includes, source_excludes	(character) fields to include in the returned document, or to exclude. a character vector
exists	(logical) Only return a logical as to whether the document exists or not.
raw	If TRUE (default), data is parsed to list. If FALSE, then raw JSON.
callopts	Curl args passed on to crul::HttpClient
verbose	If TRUE (default) the url call used printed to console.
...	Further args passed on to elastic search HTTP API as parameters.

References

<https://www.elastic.co/guide/en/elasticsearch/reference/current/docs-get.html>

Examples

```
## Not run:
(x <- connect())

if (!index_exists(x, "shakespeare")) {
  shakespeare <- system.file("examples", "shakespeare_data.json",
    package = "elastic")
  shakespeare <- type_remover(shakespeare)
  invisible(docs_bulk(x, shakespeare))
}

docs_get(x, index='shakespeare', id=10)
docs_get(x, index='shakespeare', id=12)
docs_get(x, index='shakespeare', id=12, source=TRUE)

# Get certain fields
if (gsub("\\.", "", x$ping()$version$number) < 500) {
  ### ES < v5
  docs_get(x, index='shakespeare', id=10, fields='play_name')
  docs_get(x, index='shakespeare', id=10, fields=c('play_name','speaker'))
} else {
  ### ES > v5
  docs_get(x, index='shakespeare', id=10, source='play_name')
  docs_get(x, index='shakespeare', id=10, source=c('play_name','speaker'))
}

# Just test for existence of the document
docs_get(x, index='plos', id=1, exists=TRUE)
docs_get(x, index='plos', id=123456, exists=TRUE)

# source includes / excludes
docs_get(x, index='shakespeare', id=10, source_includes = "play_name")
docs_get(x, index='shakespeare', id=10, source_excludes = "play_name")

## End(Not run)
```

docs_mget

*Get multiple documents via the multiple get API***Description**

Get multiple documents via the multiple get API

Usage

```
docs_mget(
    conn,
    index = NULL,
    type = NULL,
    ids = NULL,
    type_id = NULL,
    index_type_id = NULL,
    source = NULL,
    fields = NULL,
    raw = FALSE,
    callopts = list(),
    verbose = TRUE,
    ...
)
```

Arguments

conn	an Elasticsearch connection object, see connect()
index	Index. Required.
type	Document type. Required.
ids	More than one document id, see examples.
type_id	List of vectors of length 2, each with an element for type and id.
index_type_id	List of vectors of length 3, each with an element for index, type, and id.
source	(logical) If TRUE, return source.
fields	Fields to return from the response object.
raw	If TRUE (default), data is parsed to list. If FALSE, then raw JSON.
callopts	Curl args passed on to HttpClient
verbose	If TRUE (default) the url call used printed to console.
...	Further args passed on to elastic search HTTP API as parameters.

Details

You can pass in one of three combinations of parameters:

- Pass in something for index, type, and id. This is the simplest, allowing retrieval from the same index, same type, and many ids.
- Pass in only index and type_id - this allows you to get multiple documents from the same index, but from different types.
- Pass in only index_type_id - this is so that you can get multiple documents from different indexes and different types.

References

<https://www.elastic.co/guide/en/elasticsearch/reference/current/docs-multi-get.html>

Examples

```
## Not run:
(x <- connect())

if (!index_exists(x, 'plos')) {
  plosdat <- system.file("examples", "plos_data.json",
    package = "elastic")
  plosdat <- type_remover(plosdat)
  invisible(docs_bulk(x, plosdat))
}

# same index, many ids
docs_mget(x, index="plos", ids=c(9,10))

# Same index and type
docs_mget(x, index="plos", type="_doc", ids=c(9,10))

tmp <- docs_mget(x, index="plos", ids=c(9, 10), raw=TRUE)
es_parse(tmp)
docs_mget(x, index="plos", ids=c(9, 10), source='title')
docs_mget(x, index="plos", ids=c(14, 19), source=TRUE)

# curl options
docs_mget(x, index="plos", ids=1:2, callopts=list(verbose=TRUE))

# Same index, but different types
if (index_exists(x, 'shakespeare')) index_delete(x, 'shakespeare')
shakedat <- system.file("examples", "shakespeare_data.json",
  package = "elastic")
invisible(docs_bulk(x, shakedat))

docs_mget(x, index="shakespeare", type_id=list(c("scene",1), c("line",20)))
docs_mget(x, index="shakespeare", type_id=list(c("scene",1), c("line",20)),
  source='play_name')
```

```
# Different indices and different types pass in separately
docs_mget(x, index_type_id = list(
  c("shakespeare", "line", 20),
  c("plos", "article", 1)
)
)

## End(Not run)
```

docs_update

*Update a document***Description**

Update a document

Usage

```
docs_update(
  conn,
  index,
  id,
  body,
  type = NULL,
  fields = NULL,
  source = NULL,
  version = NULL,
  version_type = NULL,
  routing = NULL,
  parent = NULL,
  timestamp = NULL,
  ttl = NULL,
  refresh = NULL,
  timeout = NULL,
  retry_on_conflict = NULL,
  wait_for_active_shards = NULL,
  detect_noop = NULL,
  callopts = list(),
  ...
)
```

Arguments

conn	an Elasticsearch connection object, see connect()
index	(character) The name of the index. Required
id	(numeric/character) The document ID. Can be numeric or character. Required
body	The document, either a list or json

type	(character) The type of the document. optional
fields	A comma-separated list of fields to return in the response
source	Allows to control if and how the updated source should be returned in the response. By default the updated source is not returned.
version	(character) Explicit version number for concurrency control
version_type	(character) Specific version type. One of internal, external, external_gte, or force
routing	(character) Specific routing value
parent	ID of the parent document. Is only used for routing and when for the upsert request
timestamp	(date) Explicit timestamp for the document
ttl	(aka “time to live”) Expiration time for the document. Expired documents will be expunged automatically. The expiration date that will be set for a document with a provided ttl is relative to the timestamp of the document, meaning it can be based on the time of indexing or on any time provided. The provided ttl must be strictly positive and can be a number (in milliseconds) or any valid time value (e.g, 86400000, 1d).
refresh	Refresh the index after performing the operation.
timeout	(character) Explicit operation timeout, e.g., 5m (for 5 minutes)
retry_on_conflict	Specify how many times should the operation be retried when a conflict occurs (default: 0)
wait_for_active_shards	The number of shard copies required to be active before proceeding with the update operation.
detect_noop	(logical) Specifying TRUE will cause Elasticsearch to check if there are changes and, if there aren't, turn the update request into a noop.
callopts	Curl options passed on to curl::HttpClient
...	Further args to query DSL

Examples

```
## Not run:
(x <- connect())
if (!index_exists(x, 'plos')) {
  plosdat <- system.file("examples", "plos_data.json",
    package = "elastic")
  plosdat <- type_remove(plosdat)
  invisible(docs_bulk(x, plosdat))
}

docs_create(x, index='plos', id=1002,
  body=list(id="12345", title="New title"))
# and the document is there now
docs_get(x, index='plos', id=1002)
```

```
# update the document
docs_update(x, index='plos', id=1002,
  body = list(doc = list(title = "Even newer title again")))
# get it again, notice changes
docs_get(x, index='plos', id=1002)

if (!index_exists(x, 'stuffthings')) {
  index_create(x, "stuffthings")
}
docs_create(x, index='stuffthings', id=1,
  body=list(name = "foo", what = "bar"))
docs_update(x, index='stuffthings', id=1,
  body = list(doc = list(name = "hello", what = "bar")),
  source = 'name')

## End(Not run)
```

docs_update_by_query *Update documents by query*

Description

update documents by query via a POST request

Usage

```
docs_update_by_query(
  conn,
  index,
  body = NULL,
  type = NULL,
  conflicts = NULL,
  routing = NULL,
  scroll_size = NULL,
  refresh = NULL,
  wait_for_completion = NULL,
  wait_for_active_shards = NULL,
  timeout = NULL,
  scroll = NULL,
  requests_per_second = NULL,
  pipeline = NULL,
  ...
)
```

Arguments

conn	an Elasticsearch connection object, see connect()
index	(character) The name of the index. Required

body	(character/json) query to be passed on to POST request body
type	(character) The type of the document. optional
conflicts	(character) If you'd like to count version conflicts rather than cause them to abort then set conflicts=proceed
routing	(character) Specific routing value
scroll_size	(integer) By default uses scroll batches of 1000. Change batch size with this parameter.
refresh	(logical) Refresh the index after performing the operation
wait_for_completion	(logical) If wait_for_completion=FALSE then Elasticsearch will perform some preflight checks, launch the request, and then return a task which can be used with Tasks APIs to cancel or get the status of the task. Elasticsearch will also create a record of this task as a document at .tasks/task/\$taskId. This is yours to keep or remove as you see fit. When you are done with it, delete it so Elasticsearch can reclaim the space it uses. Default: TRUE
wait_for_active_shards	(logical) controls how many copies of a shard must be active before proceeding with the request.
timeout	(character) Explicit operation timeout, e.g., 5m (for 5 minutes)
scroll	(integer) control how long the "search context" is kept alive, eg scroll='10m', by default it's 5 minutes (5m)
requests_per_second	(integer) any positive decimal number (1.4, 6, 1000, etc); throttles rate at which _delete_by_query issues batches of delete operations by padding each batch with a wait time. The throttling can be disabled by setting requests_per_second=-1
pipeline	(character) a pipeline name
...	Curl args passed on to curl::verb-POST

References

<https://www.elastic.co/guide/en/elasticsearch/reference/current/docs-update-by-query.html>
<https://www.elastic.co/guide/en/elasticsearch/painless/current/painless-api-reference.html>

See Also

[docs_delete_by_query\(\)](#)

Examples

```
## Not run:
(x <- connect())
x$ping()

omdb <- system.file("examples", "omdb.json", package = "elastic")
omdb <- type_remover(omdb)
```

```
if (!index_exists(x, "omdb")) invisible(docs_bulk(x, omdb))

# can be sent without a body
docs_update_by_query(x, index='omdb')

# update
## note this works with imdbRating, a float, but didn't seem to work
## with Metascore, a long
## See link above for Painless API reference
body <- '{
  "script": {
    "source": "ctx._source.imdbRating++",
    "lang": "painless"
  },
  "query": {
    "match": {
      "Rated": "R"
    }
  }
}'
Search(x, "omdb", q = "Rated:\\\"R\\\"", asdf=TRUE,
  source = c("Title", "Rated", "imdbRating"))$hits$hits
docs_update_by_query(x, index='omdb', body = body)
Search(x, "omdb", q = "Rated:\\\"R\\\"", asdf=TRUE,
  source = c("Title", "Rated", "imdbRating"))$hits$hits

## End(Not run)
```

documents

Elasticsearch documents functions.

Description

Elasticsearch documents functions.

Details

There are five functions to work directly with documents.

- [docs_get\(\)](#)
- [docs_mget\(\)](#)
- [docs_create\(\)](#)
- [docs_delete\(\)](#)
- [docs_bulk\(\)](#)

Examples

```
## Not run:
# Get a document
# docs_get(index='plos', type='article', id=1)

# Get multiple documents
# docs_mget(index="shakespeare", type="line", id=c(9,10))

# Create a document
# docs_create(index='plos', type='article', id=35, body=list(id="12345", title="New title"))

# Delete a document
# docs_delete(index='plos', type='article', id=35)

# Bulk load documents
# plosdat <- system.file("examples", "plos_data.json", package = "elastic")
# docs_bulk(plosdat)

## End(Not run)
```

elastic

elastic

Description

An Elasticsearch R client.

About

This package gives you access to local or remote Elasticsearch databases.

Quick start

If you're connecting to a Elasticsearch server already running, skip ahead to **Search**

Install Elasticsearch (on OSX)

- Download zip or tar file from Elasticsearch see here for download: <https://www.elastic.co/downloads/elasticsearch>
- Unzip it: `untar elasticsearch-2.3.5.tar.gz`
- Move it: `sudo mv elasticsearch-2.3.5 /usr/local` (replace version with your version)
- Navigate to /usr/local: `cd /usr/local`
- Add shortcut: `sudo ln -s elasticsearch-2.3.5 elasticsearch` (replace version with your version)

For help on other platforms, see <https://www.elastic.co/guide/en/elasticsearch/reference/current/install-elasticsearch.html>

Start Elasticsearch

- Navigate to elasticsearch: `cd /usr/local/elasticsearch`
- Start elasticsearch: `bin/elasticsearch`

Initialization

The function `connect()` is used before doing anything else to set the connection details to your remote or local elasticsearch store. The details created by `connect()` are written to your options for the current session, and are used by elastic functions.

Search

The main way to search Elasticsearch is via the `Search()` function. E.g.:

`Search()`

Security

Elasticsearch is insecure out of the box! If you are running Elasticsearch locally on your own machine without exposing a port to the outside world, no worries, but if you install on a server with a public IP address, take the necessary precautions. There are a few options:

- Shield - A paid product - so probably only applicable to enterprise users
- DIY security - there are a variety of techniques for securing your Elasticsearch. I collected a number of resources in a blog post at <https://recology.info/2015/02/secure-elasticsearch/>

Elasticsearch changes

As of Elasticsearch v2:

- You can no longer create fields with dots in the name.
- Type names may not start with a dot (other than the special `.percolator` type)
- Type names may not be longer than 255 characters
- Types may no longer be deleted
- Queries and filters have been merged - all filter clauses are now query clauses. Instead, query clauses can now be used in query context or in filter context. See examples in `Search()` or `Search_uri()`

index names

The following are illegal characters, and can not be used in index names or types: `\\`, `/`, `*`, `?`, `<`, `>`, `|`, `,` (comma). double quote and whitespace are also illegal.

Author(s)

Scott Chamberlain

elastic-defunct	<i>Defunct functions in elastic</i>
-----------------	-------------------------------------

Description

- `mlt()`: The MLT API has been removed, use More Like This Query via `Search()`
- `nodes_shutdown()`: The `_shutdown` API has been removed. Instead, setup Elasticsearch to run as a service (see Running as a Service on Linux (<https://www.elastic.co/guide/en/elasticsearch/reference/2.0/setup-service.html>) or Running as a Service on Windows (<https://www.elastic.co/guide/en/elasticsearch/reference/2.0/setup-service-win.html>)) or use the `-p` command line option to write the PID to a file.
- `index_status()`: `_status` route for the index API has been removed. Replaced with the Indices Stats and Indices Recovery APIs.
- `mapping_delete()`: Elasticsearch dropped this route in their API. Instead of deleting a mapping, delete the index and recreate with a new mapping.

explain	<i>Explain a search query.</i>
---------	--------------------------------

Description

Explain a search query.

Usage

```
explain(
  conn,
  index,
  id,
  type = NULL,
  source2 = NULL,
  fields = NULL,
  routing = NULL,
  parent = NULL,
  preference = NULL,
  source = NULL,
  q = NULL,
  df = NULL,
  analyzer = NULL,
  analyze_wildcard = NULL,
  lowercase_expanded_terms = NULL,
  lenient = NULL,
  default_operator = NULL,
  source_exclude = NULL,
```

```

    source_include = NULL,
    body = NULL,
    raw = FALSE,
    ...
)

```

Arguments

conn	an Elasticsearch connection object, see connect()
index	Only one index. Required
id	Document id, only one. Required
type	Only one document type, optional
source2	(logical) Set to TRUE to retrieve the <code>_source</code> of the document explained. You can also retrieve part of the document by using <code>source_include</code> & <code>source_exclude</code> (see Get API for more details). This matches the <code>_source</code> term, but we want to avoid the leading underscore.
fields	Allows to control which stored fields to return as part of the document explained.
routing	Controls the routing in the case the routing was used during indexing.
parent	Same effect as setting the routing parameter.
preference	Controls on which shard the explain is executed.
source	Allows the data of the request to be put in the query string of the url.
q	The query string (maps to the <code>query_string</code> query).
df	The default field to use when no field prefix is defined within the query. Defaults to <code>_all</code> field.
analyzer	The analyzer name to be used when analyzing the query string. Defaults to the analyzer of the <code>_all</code> field.
analyze_wildcard	(logical) Should wildcard and prefix queries be analyzed or not. Default: FALSE
lowercase_expanded_terms	Should terms be automatically lowercased or not. Default: TRUE
lenient	If set to true will cause format based failures (like providing text to a numeric field) to be ignored. Default: FALSE
default_operator	The default operator to be used, can be AND or OR. Defaults to OR.
source_exclude	A vector of fields to exclude from the returned <code>source2</code> field
source_include	A vector of fields to extract and return from the <code>source2</code> field
body	The query definition using the Query DSL. This is passed in the body of the request.
raw	If TRUE (default), data is parsed to list. If FALSE, then raw JSON.
...	Curl args passed on to curl::HttpClient

References

<https://www.elastic.co/guide/en/elasticsearch/reference/current/search-explain.html>

Examples

```
## Not run:
(x <- connect())

explain(x, index = "plos", id = 14, q = "title:Germ")

body <- '{
  "query": {
    "match": { "title": "Germ" }
  }
}'
explain(x, index = "plos", id = 14, body=body)

## End(Not run)
```

fielddata

fielddata

Description

Deep dive on fielddata details

Details

Most fields are indexed by default, which makes them searchable. Sorting, aggregations, and accessing field values in scripts, however, requires a different access pattern from search.

Text fields use a query-time in-memory data structure called fielddata. This data structure is built on demand the first time that a field is used for aggregations, sorting, or in a script. It is built by reading the entire inverted index for each segment from disk, inverting the term-document relationship, and storing the result in memory, in the JVM heap.

fielddata is disabled on text fields by default. Fielddata can consume a lot of heap space, especially when loading high cardinality text fields. Once fielddata has been loaded into the heap, it remains there for the lifetime of the segment. Also, loading fielddata is an expensive process which can cause users to experience latency hits. This is why fielddata is disabled by default. If you try to sort, aggregate, or access values from a script on a text field, you will see this exception:

"Fielddata is disabled on text fields by default. Set fielddata=true on your_field_name in order to load fielddata in memory by uninverting the inverted index. Note that this can however use significant memory."

To enable fielddata on a text field use the PUT mapping API, for example `mapping_create("shakespeare", body = '{ "properties": { "speaker": { "type": "text", "fielddata": true } } }')`

You may get an error about `update_all_types`, in which case set `update_all_types=TRUE` in `mapping_create`, e.g.,

```
mapping_create("shakespeare", update_all_types=TRUE, body = '{ "properties": { "speaker": { "type": "text", "fielddata": true } } }')
```

See https://www.elastic.co/guide/en/elasticsearch/reference/current/fielddata.html#_enabling_fielddata_on_literal_text_literal_fields for more information.

field_caps

Field capabilities

Description

The field capabilities API allows to retrieve the capabilities of fields among multiple indices.

Usage

```
field_caps(conn, fields, index = NULL, ...)
```

Arguments

<code>conn</code>	an Elasticsearch connection object, see connect()
<code>fields</code>	A list of fields to compute stats for. required
<code>index</code>	Index name, one or more
<code>...</code>	Curl args passed on to curl::verb-GET

References

<https://www.elastic.co/guide/en/elasticsearch/reference/current/search-field-caps.html>

See Also

[field_stats\(\)](#)

Examples

```
## Not run:
x <- connect()
x$ping()

if (x$es_ver() >= 540) {
  field_caps(x, fields = "speaker", index = "shakespeare")
}

## End(Not run)
```

field_stats	<i>Search field statistics</i>
-------------	--------------------------------

Description

Search field statistics

Usage

```
field_stats(
  conn,
  fields = NULL,
  index = NULL,
  level = "cluster",
  body = list(),
  raw = FALSE,
  asdf = FALSE,
  ...
)
```

Arguments

conn	an Elasticsearch connection object, see connect()
fields	A list of fields to compute stats for. optional
index	Index name, one or more
level	Defines if field stats should be returned on a per index level or on a cluster wide level. Valid values are 'indices' and 'cluster' (default)
body	Query, either a list or json
raw	(logical) Get raw JSON back or not
asdf	(logical) If TRUE, use fromJSON to parse JSON directly to a data.frame. If FALSE (Default), list output is given.
...	Curl args passed on to curl::HttpClient

Details

The field stats api allows you to get statistical properties of a field without executing a search, but looking up measurements that are natively available in the Lucene index. This can be useful to explore a dataset which you don't know much about. For example, this allows creating a histogram aggregation with meaningful intervals based on the min/max range of values.

The field stats api by defaults executes on all indices, but can execute on specific indices too.

Note

Deprecated in Elasticsearch versions equal to/greater than 5.4.0

References

<https://www.elastic.co/guide/en/elasticsearch/reference/5.6/search-field-stats.html>

See Also

[field_caps\(\)](#)

Examples

```
## Not run:
x <- connect()

if (gsub("\\.", "", x$ping())$version$number) < 500) {
  field_stats(x, body = '{ "fields": ["speaker"] }', index = "shakespeare")
  ff <- c("scientificName", "continent", "decimalLatitude", "play_name",
    "speech_number")
  field_stats(x, "play_name")
  field_stats(x, "play_name", level = "cluster")
  field_stats(x, ff, level = "indices")
  field_stats(x, ff)
  field_stats(x, ff, index = c("gbif", "shakespeare"))

  # can also pass a body, just as with Search()
  # field_stats(x, body = list(fields = "rating")) # doesn't work
  field_stats(x, body = '{ "fields": ["scientificName"] }', index = "gbif")

  body <- '{
    "fields" : ["scientificName", "decimalLatitude"]
  }'
  field_stats(x, body = body, level = "indices", index = "gbif")
}

## End(Not run)
```

index_template

Index templates

Description

Index templates allow you to define templates that will automatically be applied when new indices are created

Usage

```
index_template_put(
  conn,
  name,
  body = NULL,
```

```

    create = NULL,
    flat_settings = NULL,
    master_timeout = NULL,
    order = NULL,
    timeout = NULL,
    ...
)

index_template_get(conn, name = NULL, filter_path = NULL, ...)

index_template_exists(conn, name, ...)

index_template_delete(conn, name, ...)

```

Arguments

conn	an Elasticsearch connection object, see connect()
name	(character) The name of the template
body	(character/list) The template definition
create	(logical) Whether the index template should only be added if new or can also replace an existing one. Default: FALSE
flat_settings	(logical) Return settings in flat format. Default: FALSE
master_timeout	(integer) Specify timeout for connection to master
order	(integer) The order for this template when merging multiple matching ones (higher numbers are merged later, overriding the lower numbers)
timeout	(integer) Explicit operation timeout
...	Curl options. Or in <code>percolate_list</code> function, further args passed on to Search()
filter_path	(character) a regex for filtering output path, see example

References

<https://www.elastic.co/guide/en/elasticsearch/reference/current/indices-templates.html>

Examples

```

## Not run:
(x <- connect())

body <- '{
  "template": "te*",
  "settings": {
    "number_of_shards": 1
  },
  "mappings": {
    "type1": {
      "_source": {

```

```

        "enabled": false
    },
    "properties": {
        "host_name": {
            "type": "keyword"
        },
        "created_at": {
            "type": "date",
            "format": "EEE MMM dd HH:mm:ss Z YYYY"
        }
    }
}
}
}'
index_template_put(x, "template_1", body = body)

# get templates
index_template_get(x)
index_template_get(x, "template_1")
index_template_get(x, c("template_1", "template_2"))
index_template_get(x, "template_*")
## filter path
index_template_get(x, "template_1", filter_path = "*.template")

# template exists
index_template_exists(x, "template_1")
index_template_exists(x, "foobar")

# delete a template
index_template_delete(x, "template_1")
index_template_exists(x, "template_1")

## End(Not run)

```

indices

Index API operations

Description

Index API operations

Usage

```

index_get(
  conn,
  index = NULL,
  features = NULL,
  raw = FALSE,
  verbose = TRUE,
  ...

```

```
)

index_exists(conn, index, ...)

index_delete(conn, index, raw = FALSE, verbose = TRUE, ...)

index_create(conn, index = NULL, body = NULL, raw = FALSE, verbose = TRUE, ...)

index_recreate(
  conn,
  index = NULL,
  body = NULL,
  raw = FALSE,
  verbose = TRUE,
  ...
)

index_close(conn, index, ...)

index_open(conn, index, ...)

index_stats(
  conn,
  index = NULL,
  metric = NULL,
  completion_fields = NULL,
  fielddata_fields = NULL,
  fields = NULL,
  groups = NULL,
  level = "indices",
  ...
)

index_settings(conn, index = "_all", ...)

index_settings_update(conn, index = NULL, body, ...)

index_segments(conn, index = NULL, ...)

index_recovery(conn, index = NULL, detailed = FALSE, active_only = FALSE, ...)

index_optimize(
  conn,
  index = NULL,
  max_num_segments = NULL,
  only_expunge_deletes = FALSE,
  flush = TRUE,
  wait_for_merge = TRUE,
```

```
    ...
)

index_forcemerger(
    conn,
    index = NULL,
    max_num_segments = NULL,
    only_expunge_deletes = FALSE,
    flush = TRUE,
    ...
)

index_upgrade(conn, index = NULL, wait_for_completion = FALSE, ...)

index_analyze(
    conn,
    text = NULL,
    field = NULL,
    index = NULL,
    analyzer = NULL,
    tokenizer = NULL,
    filters = NULL,
    char_filters = NULL,
    body = list(),
    ...
)

index_flush(
    conn,
    index = NULL,
    force = FALSE,
    full = FALSE,
    wait_if_ongoing = FALSE,
    ...
)

index_clear_cache(
    conn,
    index = NULL,
    filter = FALSE,
    filter_keys = NULL,
    fielddata = FALSE,
    query_cache = FALSE,
    id_cache = FALSE,
    ...
)

index_shrink(conn, index, index_new, body = NULL, ...)
```

Arguments

conn	an Elasticsearch connection object, see connect()
index	(character) A character vector of index names
features	(character) A single feature. One of settings, mappings, or aliases
raw	If TRUE (default), data is parsed to list. If FALSE, then raw JSON.
verbose	If TRUE (default) the url call used printed to console.
...	Curl args passed on to crul::HttpClient
body	Query, either a list or json.
metric	(character) A character vector of metrics to display. Possible values: "_all", "completion", "docs", "fielddata", "filter_cache", "flush", "get", "id_cache", "indexing", "merge", "percolate", "refresh", "search", "segments", "store", "warmer".
completion_fields	(character) A character vector of fields for completion metric (supports wildcards)
fielddata_fields	(character) A character vector of fields for fielddata metric (supports wildcards)
fields	(character) Fields to add.
groups	(character) A character vector of search groups for search statistics.
level	(character) Return stats aggregated on "cluster", "indices" (default) or "shards"
detailed	(logical) Whether to display detailed information about shard recovery. Default: FALSE
active_only	(logical) Display only those recoveries that are currently on-going. Default: FALSE
max_num_segments	(character) The number of segments the index should be merged into. Default: "dynamic"
only_expunge_deletes	(logical) Specify whether the operation should only expunge deleted documents
flush	(logical) Specify whether the index should be flushed after performing the operation. Default: TRUE
wait_for_merge	(logical) Specify whether the request should block until the merge process is finished. Default: TRUE
wait_for_completion	(logical) Should the request wait for the upgrade to complete. Default: FALSE
text	The text on which the analysis should be performed (when request body is not used)
field	Use the analyzer configured for this field (instead of passing the analyzer name)
analyzer	The name of the analyzer to use
tokenizer	The name of the tokenizer to use for the analysis
filters	A character vector of filters to use for the analysis
char_filters	A character vector of character filters to use for the analysis

force	(logical) Whether a flush should be forced even if it is not necessarily needed ie. if no changes will be committed to the index.
full	(logical) If set to TRUE a new index writer is created and settings that have been changed related to the index writer will be refreshed.
wait_if_ongoing	If TRUE, the flush operation will block until the flush can be executed if another flush operation is already executing. The default is false and will cause an exception to be thrown on the shard level if another flush operation is already running.
filter	(logical) Clear filter caches
filter_keys	(character) A vector of keys to clear when using the filter_cache parameter (default: all)
fielddata	(logical) Clear field data
query_cache	(logical) Clear query caches
id_cache	(logical) Clear ID caches for parent/child
index_new	(character) an index name, required. only applies to index_shrink method

Details

index_analyze: <https://www.elastic.co/guide/en/elasticsearch/reference/current/indices-analyze.html> This method can accept a string of text in the body, but this function passes it as a parameter in a GET request to simplify.

index_flush: <https://www.elastic.co/guide/en/elasticsearch/reference/current/indices-flush.html> From the ES website: The flush process of an index basically frees memory from the index by flushing data to the index storage and clearing the internal transaction log. By default, Elasticsearch uses memory heuristics in order to automatically trigger flush operations as required in order to clear memory.

index_status: The API endpoint for this function was deprecated in Elasticsearch v1.2.0, and will likely be removed soon. Use [index_recovery\(\)](#) instead.

index_settings_update: There are a lot of options you can change with this function. See <https://www.elastic.co/guide/en/elasticsearch/reference/current/index-modules.html#update-settings.html> for all the options.

index settings: See <https://www.elastic.co/guide/en/elasticsearch/reference/current/index-modules.html> for the *static* and *dynamic* settings you can set on indices.

Mappings

The "keyword" type is not supported in Elasticsearch < v5. If you do use a mapping with "keyword" type in Elasticsearch < v5 [index_create\(\)](#) should fail.

Author(s)

Scott Chamberlain myrmecocystus@gmail.com

References

<https://www.elastic.co/guide/en/elasticsearch/reference/current/indices.html>

Examples

```

## Not run:
# connection setup
(x <- connect())

# get information on an index
index_get(x, index='shakespeare')
## this one is the same as running index_settings('shakespeare')
index_get(x, index='shakespeare', features='settings')
index_get(x, index='shakespeare', features='mappings')
index_get(x, index='shakespeare', features='alias')

# check for index existence
index_exists(x, index='shakespeare')
index_exists(x, index='plos')

# create an index
if (index_exists(x, 'twitter')) index_delete(x, 'twitter')
index_create(x, index='twitter')
if (index_exists(x, 'things')) index_delete(x, 'things')
index_create(x, index='things')
if (index_exists(x, 'plos')) index_delete(x, 'plos')
index_create(x, index='plos')

# re-create an index
index_recreate(x, "deer")
index_recreate(x, "deer", verbose = FALSE)

# delete an index
if (index_exists(x, 'plos')) index_delete(x, index='plos')

## with a body
body <- '{
  "settings" : {
    "index" : {
      "number_of_shards" : 3,
      "number_of_replicas" : 2
    }
  }
}'
if (index_exists(x, 'alsothat')) index_delete(x, 'alsothat')
index_create(x, index='alsothat', body = body)
## with read only
body <- '{
  "settings" : {
    "index" : {
      "blocks" : {
        "read_only" : true
      }
    }
  }
}'

```

```

# index_create(x, index='myindex', body = body)
# then this delete call should fail with something like:
## > Error: 403 - blocked by: [FORBIDDEN/5/index read-only (api)]
# index_delete(x, index='myindex')

## with mappings
body <- '{
  "mappings": {
    "properties": {
      "location" : {"type" : "geo_point"}
    }
  }
}'
if (!index_exists(x, 'gbifnewgeo')) index_create(x, index='gbifnewgeo', body=body)
gbifgeo <- system.file("examples", "gbif_geosmall.json", package = "elastic")
docs_bulk(x, gbifgeo)

# close an index
index_create(x, 'plos')
index_close(x, 'plos')

# open an index
index_open(x, 'plos')

# Get stats on an index
index_stats(x, 'plos')
index_stats(x, c('plos','gbif'))
index_stats(x, c('plos','gbif'), metric='refresh')
index_stats(x, metric = "indexing")
index_stats(x, 'shakespeare', metric='completion')
index_stats(x, 'shakespeare', metric='completion', completion_fields = "completion")
index_stats(x, 'shakespeare', metric='fielddata')
index_stats(x, 'shakespeare', metric='fielddata', fielddata_fields = "evictions")
index_stats(x, 'plos', level="indices")
index_stats(x, 'plos', level="cluster")
index_stats(x, 'plos', level="shards")

# Get segments information that a Lucene index (shard level) is built with
index_segments(x)
index_segments(x, 'plos')
index_segments(x, c('plos','gbif'))

# Get recovery information that provides insight into on-going index shard recoveries
index_recovery(x)
index_recovery(x, 'plos')
index_recovery(x, c('plos','gbif'))
index_recovery(x, "plos", detailed = TRUE)
index_recovery(x, "plos", active_only = TRUE)

# Optimize an index, or many indices
if (x$es_ver() < 500) {
  ### ES < v5 - use optimize
  index_optimize(x, 'plos')
}

```

```

    index_optimize(x, c('plos','gbif'))
    index_optimize(x, 'plos')
} else {
    ### ES > v5 - use forcemerge
    index_forcemerge(x, 'plos')
}

# Upgrade one or more indices to the latest format. The upgrade process converts any
# segments written with previous formats.
if (x$es_ver() < 500) {
    index_upgrade(x, 'plos')
    index_upgrade(x, c('plos','gbif'))
}

# Performs the analysis process on a text and return the tokens breakdown
# of the text
index_analyze(x, text = 'this is a test', analyzer='standard')
index_analyze(x, text = 'this is a test', analyzer='whitespace')
index_analyze(x, text = 'this is a test', analyzer='stop')
index_analyze(x, text = 'this is a test', tokenizer='keyword',
    filters='lowercase')
index_analyze(x, text = 'this is a test', tokenizer='keyword',
    filters='lowercase', char_filters='html_strip')
index_analyze(x, text = 'this is a test', index = 'plos',
    analyzer="standard")
index_analyze(x, text = 'this is a test', index = 'shakespeare',
    analyzer="standard")

## NGram tokenizer
body <- '{
    "settings" : {
        "analysis" : {
            "analyzer" : {
                "my_ngram_analyzer" : {
                    "tokenizer" : "my_ngram_tokenizer"
                }
            },
            "tokenizer" : {
                "my_ngram_tokenizer" : {
                    "type" : "nGram",
                    "min_gram" : "2",
                    "max_gram" : "3",
                    "token_chars": [ "letter", "digit" ]
                }
            }
        }
    }
}'

if (index_exists(x, "shakespeare2")) index_delete(x, "shakespeare2")
tokenizer_set(x, index = "shakespeare2", body=body)
index_analyze(x, text = "art thouh", index = "shakespeare2",
    analyzer='my_ngram_analyzer')
```

```

# Explicitly flush one or more indices.
index_flush(x, index = "plos")
index_flush(x, index = "shakespeare")
index_flush(x, index = c("plos", "shakespeare"))
index_flush(x, index = "plos", wait_if_ongoing = TRUE)
index_flush(x, index = "plos", verbose = TRUE)

# Clear either all caches or specific cached associated with one ore more indices.
index_clear_cache(x)
index_clear_cache(x, index = "plos")
index_clear_cache(x, index = "shakespeare")
index_clear_cache(x, index = c("plos", "shakespeare"))
index_clear_cache(x, filter = TRUE)

# Index settings
## get settings
index_settings(x)
index_settings(x, "_all")
index_settings(x, 'gbif')
index_settings(x, c('gbif', 'plos'))
index_settings(x, '*s')
## update settings
if (index_exists(x, 'foobar')) index_delete(x, 'foobar')
index_create(x, "foobar")
settings <- list(index = list(number_of_replicas = 4))
index_settings_update(x, "foobar", body = settings)
index_get(x, "foobar")$foobar$settings

# Shrink index - Can only shrink an index if it has >1 shard
## index must be read only, a copy of every shard in the index must
## reside on the same node, and the cluster health status must be green
### index_settings_update call to change these
settings <- list(
  index.routing.allocation.require._name = "shrink_node_name",
  index.blocks.write = "true"
)
if (index_exists(x, 'barbarbar')) index_delete(x, 'barbarbar')
index_create(x, "barbarbar")
index_settings_update(x, "barbarbar", body = settings)
cat_recovery(x, index='barbarbar')
# index_shrink(x, "barbarbar", "barfoobbar")

## End(Not run)

```

ingest

Ingest API operations

Description

Ingest API operations

Usage

```
pipeline_create(conn, id, body, ...)

pipeline_attachment(conn, index, id, pipeline, body, type = NULL, ...)

pipeline_get(conn, id, filter_path = NULL, ...)

pipeline_delete(conn, id, body, ...)

pipeline_simulate(conn, body, id = NULL, ...)
```

Arguments

conn	an Elasticsearch connection object, see connect()
id	(character) one or more pipeline id's. with delete, you can use a wildcard match
body	body describing pipeline, see examples and Elasticsearch docs
...	Curl args passed on to crul::verb-POST , crul::verb-GET , crul::verb-PUT , or crul::verb-DELETE
index	(character) an index. only used in pipeline_attachment
pipeline	(character) a pipeline name. only used in pipeline_attachment
type	(character) a type. only used in pipeline_attachment. by default ths is set to NULL - optional in ES <= v6.3; not allowed in ES >= v6.4
filter_path	(character) fields to return. deafults to all if not given

Details

ingest/pipeline functions available in Elasticsearch v5 and greater

Value

a named list

Attachments

See <https://www.elastic.co/guide/en/elasticsearch/plugins/current/ingest-attachment.html> You need to install the attachment processor plugin to be able to use attachments in pipelines

References

<https://www.elastic.co/guide/en/elasticsearch/reference/current/ingest-apis.html>,
<https://www.elastic.co/guide/en/elasticsearch/plugins/current/using-ingest-attachment.html>

Examples

```
## Not run:
# connection setup
(x <- connect())

# create
body1 <- '{
  "description" : "do a thing",
  "version" : 123,
  "processors" : [
    {
      "set" : {
        "field": "foo",
        "value": "bar"
      }
    }
  ]
}'
body2 <- '{
  "description" : "do another thing",
  "processors" : [
    {
      "set" : {
        "field": "stuff",
        "value": "things"
      }
    }
  ]
}'
pipeline_create(x, id = 'foo', body = body1)
pipeline_create(x, id = 'bar', body = body2)

# get
pipeline_get(x, id = 'foo')
pipeline_get(x, id = 'bar')
pipeline_get(x, id = 'foo', filter_path = "*.version")
pipeline_get(x, id = c('foo', 'bar')) # get >1

# delete
pipeline_delete(x, id = 'foo')

# simulate
## with pipeline included
body <- '{
  "pipeline" : {
    "description" : "do another thing",
    "processors" : [
      {
        "set" : {
          "field": "stuff",
          "value": "things"
        }
      }
    ]
  }
}
```

```

    }
  ]
},
"docs" : [
  { "_source": {"foo": "bar"} },
  { "_source": {"foo": "world"} }
]
}'
pipeline_simulate(x, body)

## referencing existing pipeline
body <- '{
  "docs" : [
    { "_source": {"foo": "bar"} },
    { "_source": {"foo": "world"} }
  ]
}'
pipeline_simulate(x, body, id = "foo")

# attachments - Note: you need the attachment plugin for this, see above
body1 <- '{
  "description" : "do a thing",
  "version" : 123,
  "processors" : [
    {
      "attachment" : {
        "field" : "data"
      }
    }
  ]
}'
pipeline_create(x, "baz", body1)
body_attach <- '{
  "data": "e1xydGYxXGFuc2kNCkxvcnVtIGlwc3VtIGRvbG9yIHNPdCBhbWV0DQpccGFyIH0="
}'
if (!index_exists(x, "boomarang")) index_create(x, "boomarang")
docs_create(x, 'boomarang', id = 1, body = list(title = "New title"))
pipeline_attachment(x, "boomarang", "1", "baz", body_attach)
pipeline_get(x, id = 'baz')

## End(Not run)

```

Description

Mapping management

Usage

```

mapping_create(
  conn,
  index,
  body,
  type = NULL,
  update_all_types = FALSE,
  include_type_name = NULL,
  ...
)

mapping_get(conn, index = NULL, type = NULL, include_type_name = NULL, ...)

field_mapping_get(
  conn,
  index = NULL,
  type = NULL,
  field,
  include_defaults = FALSE,
  include_type_name = NULL,
  ...
)

type_exists(conn, index, type, ...)

```

Arguments

conn	an Elasticsearch connection object, see connect()
index	(character) An index
body	(list) Either a list or json, representing the query.
type	(character) A document type
update_all_types	(logical) update all types. default: FALSE. This parameter is deprecated in ES v6.3.0 and higher, see https://github.com/elastic/elasticsearch/pull/28284
include_type_name	(logical) If set to TRUE, you can include a type name, if not an error will occur. default: not set. See Details.
...	Curl options passed on to curl::verb-PUT , curl::verb-GET , or curl::verb-HEAD
field	(character) One or more field names
include_defaults	(logical) Whether to return default values

Details

Find documentation for each function at:

- `mapping_create` - <https://www.elastic.co/guide/en/elasticsearch/reference/current/indices-put-mapping.html>
- `type_exists` - <https://www.elastic.co/guide/en/elasticsearch/reference/current/indices-types-exists.html>
- `mapping_delete` - FUNCTION DEFUNCT - instead of deleting mapping, delete index and recreate index with new mapping
- `mapping_get` - <https://www.elastic.co/guide/en/elasticsearch/reference/current/indices-get-mapping.html>
- `field_mapping_get` - <https://www.elastic.co/guide/en/elasticsearch/reference/current/indices-get-field-mapping.html>

See <https://www.elastic.co/guide/en/elasticsearch/reference/current/removal-of-types.html> for information on type removal

Examples

```
## Not run:
# connection setup
(x <- connect())

# Used to check if a type/types exists in an index/indices
type_exists(x, index = "plos", type = "article")
type_exists(x, index = "plos", type = "articles")
type_exists(x, index = "shakespeare", type = "line")

# The put mapping API allows to register specific mapping definition for a specific type.
## a good mapping body
body <- list(properties = list(
  journal = list(type="text"),
  year = list(type="long")
))
if (!index_exists(x, "plos")) index_create(x, "plos")
mapping_create(x, index = "plos", type = "citation", body=body)
## OR if above fails, try
mapping_create(x, index = "plos", type = "citation", body=body,
  include_type_name=TRUE)
## ES >= 7, no type
mapping_create(x, index = "plos", body=body)

### or as json
body <- '{
  "properties": {
    "journal": { "type": "text" },
    "year": { "type": "long" }
  }
}'
mapping_create(x, index = "plos", type = "citation", body=body)
mapping_get(x, "plos", "citation")

## A bad mapping body
body <- list(things = list(properties = list(
  journal = list("text")
```

```

)))
# mapping_create(x, index = "plos", type = "things", body=body)

# Get mappings
mapping_get(x, '_all')
mapping_get(x, index = "plos")
mapping_get(x, index = c("shakespeare","plos"))
# mapping_get(x, index = "shakespeare", type = "act")
# mapping_get(x, index = "shakespeare", type = c("act","line"))

# Get field mappings
plosdat <- system.file("examples", "plos_data.json",
  package = "elastic")
plosdat <- type_remover(plosdat)
invisible(docs_bulk(x, plosdat))
field_mapping_get(x, index = "_all", field = "text")
field_mapping_get(x, index = "plos", field = "title")
field_mapping_get(x, index = "plos", field = "*")
field_mapping_get(x, index = "plos", field = "title", include_defaults = TRUE)
field_mapping_get(x, type = c("article","record"), field = c("title","class"))
field_mapping_get(x, type = "a*", field = "t*")

# Create geospatial mapping
if (index_exists(x, "gbifgeopoint")) index_delete(x, "gbifgeopoint")
file <- system.file("examples", "gbif_geopoint.json",
  package = "elastic")
file <- type_remover(file)
index_create(x, "gbifgeopoint")
body <- '{
  "properties" : {
    "location" : { "type" : "geo_point" }
  }
}'
mapping_create(x, "gbifgeopoint", body = body)
invisible(docs_bulk(x, file))

# update_all_fields, see also ?fielddata
if (x$es_ver() < 603) {
  mapping_create(x, "shakespeare", "record", update_all_types=TRUE, body = '{
    "properties": {
      "speaker": {
        "type": "text",
        "fielddata": true
      }
    }
  }')
} else {
  index_create(x, 'brownchair')
  mapping_create(x, 'brownchair', body = '{
    "properties": {
      "foo": {
        "type": "text",
        "fielddata": true
      }
    }
  }')
}

```

```

    }
  }
}'')
}

## End(Not run)

```

msearch

Multi-search

Description

Performs multiple searches, defined in a file

Usage

```
msearch(conn, x, raw = FALSE, asdf = FALSE, ...)
```

Arguments

conn	an Elasticsearch connection object, see connect()
x	(character) A file path
raw	(logical) Get raw JSON back or not.
asdf	(logical) If TRUE, use jsonlite::fromJSON() to parse JSON directly to a data.frame. If FALSE (Default), list output is given.
...	Curl args passed on to crul::verb-POST

Details

This function behaves similarly to [docs_bulk\(\)](#) - performs searches based on queries defined in a file.

See Also

[Search_uri\(\)](#) [Search\(\)](#)

Examples

```

## Not run:
x <- connect()

msearch1 <- system.file("examples", "msearch_eg1.json", package = "elastic")
readLines(msearch1)
msearch(x, msearch1)

tf <- tempfile(fileext = ".json")
cat('{"index" : "shakespeare"}', file = tf, sep = "\n")

```

```

cat('{"query" : {"match_all" : {}}, "from" : 0, "size" : 5}', sep = "\n",
    file = tf, append = TRUE)
readLines(tf)
msearch(x, tf)

## End(Not run)

```

mtermvectors

Multi Termvectors

Description

Multi Termvectors

Usage

```

mtermvectors(
  conn,
  index = NULL,
  type = NULL,
  ids = NULL,
  body = list(),
  pretty = TRUE,
  field_statistics = TRUE,
  fields = NULL,
  offsets = TRUE,
  parent = NULL,
  payloads = TRUE,
  positions = TRUE,
  preference = "random",
  realtime = TRUE,
  routing = NULL,
  term_statistics = FALSE,
  version = NULL,
  version_type = NULL,
  ...
)

```

Arguments

conn	an Elasticsearch connection object, see connect()
index	(character) The index in which the document resides.
type	(character) The type of the document.
ids	(character) One or more document ids
body	(character) Define parameters and or supply a document to get termvectors for
pretty	(logical) pretty print. Default: TRUE

field_statistics	(character) Specifies if document count, sum of document frequencies and sum of total term frequencies should be returned. Default: TRUE
fields	(character) A comma-separated list of fields to return.
offsets	(character) Specifies if term offsets should be returned. Default: TRUE
parent	(character) Parent id of documents.
payloads	(character) Specifies if term payloads should be returned. Default: TRUE
positions	(character) Specifies if term positions should be returned. Default: TRUE
preference	(character) Specify the node or shard the operation should be performed on (Default: random).
realtime	(character) Specifies if request is real-time as opposed to near-real-time (Default: TRUE).
routing	(character) Specific routing value.
term_statistics	(character) Specifies if total term frequency and document frequency should be returned. Default: FALSE
version	(character) Explicit version number for concurrency control
version_type	(character) Specific version type, valid choices are: 'internal', 'external', 'external_gte', 'force'
...	Curl args passed on to curl::verb-POST

Details

Multi termvectors API allows to get multiple termvectors based on an index, type and id.

References

<https://www.elastic.co/guide/en/elasticsearch/reference/current/docs-multi-termvectors.html>

See Also

[termvectors\(\)](#)

Examples

```
## Not run:
x <- connect()

if (index_exists(x, 'omdb')) index_delete(x, "omdb")
omdb <- system.file("examples", "omdb.json", package = "elastic")
omdb <- type_remover(omdb)
invisible(docs_bulk(x, omdb))
out <- Search(x, "omdb", size = 2)$hits$hits
ids <- vapply(out, "[[", "", "_id")

# no index
```

```

body <- '{
  "docs": [
    {
      "_index": "omdb",
      "_id": "%s",
      "term_statistics": true
    },
    {
      "_index": "omdb",
      "_id": "%s",
      "fields": [
        "Plot"
      ]
    }
  ]
}'
mtermvectors(x, body = sprintf(body, ids[1], ids[2]))

# index given
body <- '{
  "docs": [
    {
      "_id": "%s",
      "fields": [
        "Plot"
      ],
      "term_statistics": true
    },
    {
      "_id": "%s",
      "fields": [
        "Title"
      ]
    }
  ]
}'
mtermvectors(x, 'omdb', body = sprintf(body, ids[1], ids[2]))

# parameters same for both documents, so can simplify
body <- '{
  "ids" : ["%s", "%s"],
  "parameters": {
    "fields": [
      "Plot"
    ],
    "term_statistics": true
  }
}'
mtermvectors(x, 'omdb', body = sprintf(body, ids[1], ids[2]))

# you can give user provided documents via the 'docs' parameter
## though you have to give index and type that exist in your Elasticsearch
## instance

```

```

body <- '{
  "docs": [
    {
      "_index": "omdb",
      "doc" : {
        "Director" : "John Doe",
        "Plot" : "twitter test test test"
      }
    },
    {
      "_index": "omdb",
      "doc" : {
        "Director" : "Jane Doe",
        "Plot" : "Another twitter test ..."
      }
    }
  ]
}'
mtermvectors(x, body = body)

## End(Not run)

```

nodes	<i>Elasticsearch nodes endpoints.</i>
-------	---------------------------------------

Description

Elasticsearch nodes endpoints.

Usage

```
nodes_stats(conn, node = NULL, metric = NULL, raw = FALSE, fields = NULL, ...)
```

```
nodes_info(conn, node = NULL, metric = NULL, raw = FALSE, ...)
```

```

nodes_hot_threads(
  conn,
  node = NULL,
  metric = NULL,
  threads = 3,
  interval = "500ms",
  type = NULL,
  raw = FALSE,
  ...
)

```

Arguments

<code>conn</code>	an Elasticsearch connection object, see connect()
<code>node</code>	The node
<code>metric</code>	A metric to get. See Details.
<code>raw</code>	If TRUE (default), data is parsed to list. If FALSE, then raw JSON.
<code>fields</code>	You can get information about field data memory usage on node level or on index level
<code>...</code>	Curl args passed on to curl::verb-GET
<code>threads</code>	(character) Number of hot threads to provide. Default: 3
<code>interval</code>	(character) The interval to do the second sampling of threads. Default: 500ms
<code>type</code>	(character) The type to sample, defaults to <code>cpu</code> , but supports <code>wait</code> and <code>block</code> to see hot threads that are in wait or block state.

Details

<https://www.elastic.co/guide/en/elasticsearch/reference/current/cluster-nodes-stats.html>

By default, all stats are returned. You can limit this by combining any of `indices`, `os`, `process`, `jvm`, `network`, `transport`, `http`, `fs`, `breaker` and `thread_pool`. With the `metric` parameter you can select zero or more of:

- `indices` Indices stats about size, document count, indexing and deletion times, search times, field cache size, merges and flushes
- `os` retrieve information that concern the operating system
- `fs` File system information, data path, free disk space, read/write stats
- `http` HTTP connection information
- `jvm` JVM stats, memory pool information, garbage collection, buffer pools
- `network` TCP information
- `os` Operating system stats, load average, cpu, mem, swap
- `process` Process statistics, memory consumption, cpu usage, open file descriptors
- `thread_pool` Statistics about each thread pool, including current size, queue and rejected tasks
- `transport` Transport statistics about sent and received bytes in cluster communication
- `breaker` Statistics about the field data circuit breaker

[nodes_hot_threads\(\)](#) returns plain text, so `base::cat()` is used to print to the console.

Examples

```
## Not run:
# connection setup
(x <- connect())

(out <- nodes_stats(x))
nodes_stats(x, node = names(out$nodes))
```

```

nodes_stats(x, metric='get')
nodes_stats(x, metric='jvm')
nodes_stats(x, metric=c('os','process'))
nodes_info(x)
nodes_info(x, metric='process')
nodes_info(x, metric='jvm')
nodes_info(x, metric='http')
nodes_info(x, metric='network')

## End(Not run)

```

percolate

Percolater

Description

Store queries into an index then, via the percolate API, define documents to retrieve these queries.

Usage

```

percolate_register(
  conn,
  index,
  id,
  type = NULL,
  body = list(),
  routing = NULL,
  preference = NULL,
  ignore_unavailable = NULL,
  percolate_format = NULL,
  refresh = NULL,
  ...
)

percolate_match(
  conn,
  index,
  type = NULL,
  body,
  routing = NULL,
  preference = NULL,
  ignore_unavailable = NULL,
  percolate_format = NULL,
  ...
)

percolate_list(conn, index, ...)

```

```
percolate_count(conn, index, type, body, ...)
```

```
percolate_delete(conn, index, id, ...)
```

Arguments

conn	an Elasticsearch connection object, see connect()
index	Index name. Required
id	A precolator id. Required
type	Document type. Required
body	Body json, or R list.
routing	(character) In case the percolate queries are partitioned by a custom routing value, that routing option makes sure that the percolate request only gets executed on the shard where the routing value is partitioned to. This means that the percolate request only gets executed on one shard instead of all shards. Multiple values can be specified as a comma separated string, in that case the request can be executed on more than one shard.
preference	(character) Controls which shard replicas are preferred to execute the request on. Works the same as in the search API.
ignore_unavailable	(logical) Controls if missing concrete indices should silently be ignored. Same as is in the search API.
percolate_format	(character) If ids is specified then the matches array in the percolate response will contain a string array of the matching ids instead of an array of objects. This can be useful to reduce the amount of data being send back to the client. Obviously if there are two percolator queries with same id from different indices there is no way to find out which percolator query belongs to what index. Any other value to percolate_format will be ignored.
refresh	If TRUE then refresh the affected shards to make this operation visible to search, if "wait_for" then wait for a refresh to make this operation visible to search, if FALSE (default) then do nothing with refreshes. Valid choices: TRUE, FALSE, "wait_for"
...	Curl options. Or in percolate_list function, further args passed on to Search()

Details

Additional body options, pass those in the body. These aren't query string parameters:

- filter - Reduces the number queries to execute during percolating. Only the percolator queries that match with the filter will be included in the percolate execution. The filter option works in near realtime, so a refresh needs to have occurred for the filter to included the latest percolate queries.
- query - Same as the filter option, but also the score is computed. The computed scores can then be used by the track_scores and sort option.
- size - Defines to maximum number of matches (percolate queries) to be returned. Defaults to unlimited.

- `track_scores` - Whether the `_score` is included for each match. The `_score` is based on the query and represents how the query matched the percolate query's metadata, not how the document (that is being percolated) matched the query. The `query` option is required for this option. Defaults to `false`.
- `sort` - Define a sort specification like in the search API. Currently only sorting `_score` reverse (default relevancy) is supported. Other sort fields will throw an exception. The `size` and `query` option are required for this setting. Like `track_score` the score is based on the query and represents how the query matched to the percolate query's metadata and not how the document being percolated matched to the query.
- `aggs` - Allows aggregation definitions to be included. The aggregations are based on the matching percolator queries, look at the aggregation documentation on how to define aggregations.
- `highlight` - Allows highlight definitions to be included. The document being percolated is being highlight for each matching query. This allows you to see how each match is highlighting the document being percolated. See highlight documentation on how to define highlights. The `size` option is required for highlighting, the performance of highlighting in the percolate API depends of how many matches are being highlighted.

The Elasticsearch v5 split

In Elasticsearch < v5, there's a certain set of percolate APIs available, while in Elasticsearch >= v5, there's a different set of APIs available.

Internally within these percolate functions we detect your Elasticsearch version, then use the appropriate APIs

References

<https://www.elastic.co/guide/en/elasticsearch/reference/current/query-dsl-percolate-query.html>

Examples

```
## Not run:
x <- connect(errors = "complete")

##### Elasticsearch < v5
if (x$es_ver() < 500) {
  # typical usage
  ## create an index first
  if (index_exists(x, "myindex")) index_delete(x, "myindex")
  mapping <- '{
    "mappings": {
      "mytype": {
        "properties": {
          "message": {
            "type": "text"
          },
          "query": {
            "type": "percolator"
          }
        }
      }
    }
  }
```

```

    }
  }
}
}'
index_create(x, "myindex", body = mapping)

## register a percolator
perc_body = '{
  "query" : {
    "match" : {
      "message" : "bonsai tree"
    }
  }
}'
percolate_register(x, index = "myindex", type = "mytype",
  id = 1, body = perc_body)

## register another
perc_body2 <- '{
  "query" : {
    "match" : {
      "message" : "jane doe"
    }
  }
}'
percolate_register(x, index = "myindex", type = "mytype",
  id = 2, body = perc_body2)

## match a document to a percolator
doc <- '{
  "query": {
    "percolate": {
      "field": "query",
      "document": {
        "message" : "A new bonsai tree in the office"
      }
    }
  }
}'
percolate_match(x, index = "myindex", type = "mytype", body = doc)

## List percolators - for an index, no type, can't do across indices
percolate_list(x, index = "myindex")$hits$hits

## Percolate counter
percolate_count(x, index = "myindex", type = "mytype", body = doc)$total

## delete a percolator
percolate_delete(x, index = "myindex", id = 2)
} # end ES < 5

##### Elasticsearch >= v5

```

```

if (x$es_ver() >= 500 && x$es_ver() <= 700) {
  if (index_exists(x, "myindex")) index_delete(x, "myindex")

  body <- '{
    "mappings": {
      "mytype": {
        "properties": {
          "message": {
            "type": "text"
          },
          "query": {
            "type": "percolator"
          }
        }
      }
    }
  }'

  # create the index with mapping
  index_create(x, "myindex", body = body)

  ## register a percolator
  z <- '{
    "query" : {
      "match" : {
        "message" : "bonsai tree"
      }
    }
  }'
  percolate_register(x, index = "myindex", type = "mytype", id = 1, body = z)

  ## register another
  x2 <- '{
    "query" : {
      "match" : {
        "message" : "the office"
      }
    }
  }'
  percolate_register(x, index = "myindex", type = "mytype", id = 2, body = x2)

  ## match a document to a percolator
  query <- '{
    "query" : {
      "percolate" : {
        "field": "query",
        "document": {
          "message": "A new bonsai tree in the office"
        }
      }
    }
  }'
  percolate_match(x, index = "myindex", body = query)

```

```

} # end ES >= 5

##### Elasticsearch >= v7
if (x$es_ver() >= 700) {
  if (index_exists(x, "myindex")) index_delete(x, "myindex")

  body <- '{
    "mappings": {
      "properties": {
        "message": {
          "type": "text"
        },
        "query": {
          "type": "percolator"
        }
      }
    }
  }'

  # create the index with mapping
  index_create(x, "myindex", body = body)

  ## register a percolator
  z <- '{
    "query" : {
      "match" : {
        "message" : "bonsai tree"
      }
    }
  }'
  percolate_register(x, index = "myindex", id = 1, body = z)

  ## register another
  x2 <- '{
    "query" : {
      "match" : {
        "message" : "the office"
      }
    }
  }'
  percolate_register(x, index = "myindex", id = 2, body = x2)

  ## match a document to a percolator
  query <- '{
    "query" : {
      "percolate" : {
        "field": "query",
        "document": {
          "message": "A new bonsai tree in the office"
        }
      }
    }
  }'

```

```
    }  
  }  
}'  
percolate_match(x, index = "myindex", body = query)  
} # end ES >= 7
```

```
## End(Not run)
```

ping

Ping an Elasticsearch server.

Description

Ping an Elasticsearch server.

Usage

```
ping(conn, ...)
```

Arguments

conn	an Elasticsearch connection object, see connect()
...	Curl args passed on to crul::verb-GET

See Also

[connect\(\)](#)

Examples

```
## Not run:  
x <- connect()  
ping(x)  
# ideally call ping on the connetion object itself  
x$ping()  
  
## End(Not run)
```

preference	<i>Preferences.</i>
------------	---------------------

Description

Preferences.

Details

- `_primary` The operation will go and be executed only on the primary shards.
- `_primary_first` The operation will go and be executed on the primary shard, and if not available (failover), will execute on other shards.
- `_local` The operation will prefer to be executed on a local allocated shard if possible.
- `_only_node:xyz` Restricts the search to execute only on a node with the provided node id (xyz in this case).
- `_prefer_node:xyz` Prefers execution on the node with the provided node id (xyz in this case) if applicable.
- `_shards:2,3` Restricts the operation to the specified shards. (2 and 3 in this case). This preference can be combined with other preferences but it has to appear first: `_shards:2,3;_primary`
- Custom (string) value A custom value will be used to guarantee that the same shards will be used for the same custom value. This can help with "jumping values" when hitting different shards in different refresh states. A sample value can be something like the web session id, or the user name.

reindex	<i>Reindex</i>
---------	----------------

Description

Reindex all documents from one index to another.

Usage

```
reindex(
  conn,
  body,
  refresh = NULL,
  requests_per_second = NULL,
  slices = NULL,
  timeout = NULL,
  wait_for_active_shards = NULL,
  wait_for_completion = NULL,
  ...
)
```

Arguments

<code>conn</code>	an Elasticsearch connection object, see connect()
<code>body</code>	(list/character/json) The search definition using the Query DSL and the prototype for the index request.
<code>refresh</code>	(logical) Should the effected indexes be refreshed?
<code>requests_per_second</code>	(integer) The throttle to set on this request in sub-requests per second. - 1 means no throttle. Default: 0
<code>slices</code>	(integer) The number of slices this task should be divided into. Defaults to 1 meaning the task isn't sliced into subtasks. Default: 1
<code>timeout</code>	(character) Time each individual bulk request should wait for shards that are unavailable. Default: '1m'
<code>wait_for_active_shards</code>	(integer) Sets the number of shard copies that must be active before proceeding with the reindex operation. Defaults to 1, meaning the primary shard only. Set to all for all shard copies, otherwise set to any non-negative value less than or equal to the total number of copies for the shard (number of replicas + 1)
<code>wait_for_completion</code>	(logical) Should the request block until the reindex is complete? Default: TRUE
<code>...</code>	Curl options, passed on to curl::verb-POST

References

<https://www.elastic.co/guide/en/elasticsearch/reference/current/docs-reindex.html>

Examples

```
## Not run:
x <- connect()

if (!index_exists(x, "twitter")) index_create(x, "twitter")
if (!index_exists(x, "new_twitter")) index_create(x, "new_twitter")
body <- '{
  "source": {
    "index": "twitter"
  },
  "dest": {
    "index": "new_twitter"
  }
}'
reindex(x, body = body)

## End(Not run)
```

scroll	<i>Scroll search function</i>
--------	-------------------------------

Description

Scroll search function

Usage

```
scroll(
  conn,
  x,
  time_scroll = "1m",
  raw = FALSE,
  asdf = FALSE,
  stream_opts = list(),
  ...
)
```

```
scroll_clear(conn, x = NULL, all = FALSE, ...)
```

Arguments

conn	an Elasticsearch connection object, see connect()
x	(character) For scroll, a single scroll id; for scroll_clear, one or more scroll id's
time_scroll	(character) Specify how long a consistent view of the index should be maintained for scrolled search, e.g., "30s", "1m". See units-time .
raw	(logical) If FALSE (default), data is parsed to list. If TRUE, then raw JSON.
asdf	(logical) If TRUE, use jsonlite::fromJSON() to parse JSON directly to a data.frame. If FALSE (Default), list output is given.
stream_opts	(list) A list of options passed to jsonlite::stream_out() - Except that you can't pass x as that's the data that's streamed out, and pass a file path instead of a connection to con. pagesize param doesn't do much as that's more or less controlled by paging with ES.
...	Curl args passed on to crul::verb-POST
all	(logical) If TRUE (default) then all search contexts cleared. If FALSE, scroll id's must be passed to x

Value

scroll() returns a list, identical to what [Search\(\)](#) returns. With attribute scroll that is the scroll value set via the time_scroll parameter

scroll_clear() returns a boolean (TRUE on success)

Scores

Scores will be the same for all documents that are returned from a scroll request. Dents da rules.

Inputs

Inputs to `scroll()` can be one of:

- list - This usually will be the output of `Search()`, but you could in theory make a list yourself with the appropriate elements
- character - A scroll ID - this is typically the scroll id output from a call to `Search()`, accessed like `res$`_scroll_id``

All other classes passed to `scroll()` will fail with message

Lists passed to `scroll()` without a `_scroll_id` element will trigger an error.

From lists output from `Search()` there should be an attribute ("scroll") that is the scroll value set in the `Search()` request - if that attribute is missing from the list, we'll attempt to use the `time_scroll` parameter value set in the `scroll()` function call

The output of `scroll()` has the scroll time value as an attribute so the output can be passed back into `scroll()` to continue.

Clear scroll

Search context are automatically removed when the scroll timeout has been exceeded. Keeping scrolls open has a cost, so scrolls should be explicitly cleared as soon as the scroll is not being used anymore using `scroll_clear`

Sliced scrolling

For scroll queries that return a lot of documents it is possible to split the scroll in multiple slices which can be consumed independently.

See the example in this man file.

Aggregations

If the request specifies aggregations, only the initial search response will contain the aggregations results.

See Also

[Search\(\)](#)

Examples

```
## Not run:
# connection setup
(con <- connect())

# Basic usage - can use across all indices
res <- Search(con, time_scroll="1m")
```

```

scroll(con, res)$`_scroll_id`

# use on a specific index - and specify a query
res <- Search(con, index = 'shakespeare', q="a*", time_scroll="1m")
res$`_scroll_id`

# Setting "sort=_doc" to turn off sorting of results - faster
res <- Search(con, index = 'shakespeare', q="a*", time_scroll="1m",
  body = '{"sort": ["_doc"]}')
res$`_scroll_id`

# Pass scroll_id to scroll function
scroll(con, res$`_scroll_id`)

# Get all results - one approach is to use a while loop
res <- Search(con, index = 'shakespeare', q="a*", time_scroll="5m",
  body = '{"sort": ["_doc"]}')
out <- res$hits$hits
hits <- 1
while(hits != 0){
  res <- scroll(con, res$`_scroll_id`, time_scroll="5m")
  hits <- length(res$hits$hits)
  if(hits > 0)
    out <- c(out, res$hits$hits)
}
length(out)
res$hits$total
out[[1]]

# clear scroll
## individual scroll id
res <- Search(con, index = 'shakespeare', q="a*", time_scroll="5m",
  body = '{"sort": ["_doc"]}')
scroll_clear(con, res$`_scroll_id`)

## many scroll ids
res1 <- Search(con, index = 'shakespeare', q="c*", time_scroll="5m",
  body = '{"sort": ["_doc"]}')
res2 <- Search(con, index = 'shakespeare', q="d*", time_scroll="5m",
  body = '{"sort": ["_doc"]}')
nodes_stats(con, metric = "indices")$nodes[[1]]$indices$search$open_contexts
scroll_clear(con, c(res1$`_scroll_id`, res2$`_scroll_id`))
nodes_stats(con, metric = "indices")$nodes[[1]]$indices$search$open_contexts

## all scroll ids
res1 <- Search(con, index = 'shakespeare', q="f*", time_scroll="1m",
  body = '{"sort": ["_doc"]}')
res2 <- Search(con, index = 'shakespeare', q="g*", time_scroll="1m",
  body = '{"sort": ["_doc"]}')
res3 <- Search(con, index = 'shakespeare', q="k*", time_scroll="1m",
  body = '{"sort": ["_doc"]}')
scroll_clear(con, all = TRUE)

```

```

## sliced scrolling
body1 <- '{
  "slice": {
    "id": 0,
    "max": 2
  },
  "query": {
    "match" : {
      "text_entry" : "a*"
    }
  }
}'

body2 <- '{
  "slice": {
    "id": 1,
    "max": 2
  },
  "query": {
    "match" : {
      "text_entry" : "a*"
    }
  }
}'

res1 <- Search(con, index = 'shakespeare', time_scroll="1m", body = body1)
res2 <- Search(con, index = 'shakespeare', time_scroll="1m", body = body2)
scroll(con, res1$`_scroll_id`)
scroll(con, res2$`_scroll_id`)

out1 <- list()
hits <- 1
while(hits != 0){
  tmp1 <- scroll(con, res1$`_scroll_id`)
  hits <- length(tmp1$hits$hits)
  if(hits > 0)
    out1 <- c(out1, tmp1$hits$hits)
}

out2 <- list()
hits <- 1
while(hits != 0){
  tmp2 <- scroll(con, res2$`_scroll_id`)
  hits <- length(tmp2$hits$hits)
  if(hits > 0)
    out2 <- c(out2, tmp2$hits$hits)
}

c(
  lapply(out1, "[", "_source"),
  lapply(out2, "[", "_source")
)

```

```

# using jsonlite::stream_out
res <- Search(con, time_scroll = "1m")
file <- tempfile()
scroll(con,
  x = res$`_scroll_id`,
  stream_opts = list(file = file)
)
jsonlite::stream_in(file(file))
unlink(file)

## stream_out and while loop
(file <- tempfile())
res <- Search(con, index = "shakespeare", time_scroll = "5m",
  size = 1000, stream_opts = list(file = file))
while(!inherits(res, "warning")) {
  res <- tryCatch(scroll(
    conn = con,
    x = res$`_scroll_id`,
    time_scroll = "5m",
    stream_opts = list(file = file)
  ), warning = function(w) w)
}
NROW(df <- jsonlite::stream_in(file(file)))
head(df)

## End(Not run)

```

Search

Full text search of Elasticsearch

Description

Full text search of Elasticsearch

Usage

```

Search(
  conn,
  index = NULL,
  type = NULL,
  q = NULL,
  df = NULL,
  analyzer = NULL,
  default_operator = NULL,
  explain = NULL,
  source = NULL,
  fields = NULL,
  sort = NULL,

```

```

track_scores = NULL,
timeout = NULL,
terminate_after = NULL,
from = NULL,
size = NULL,
search_type = NULL,
lowercase_expanded_terms = NULL,
analyze_wildcard = NULL,
version = NULL,
lenient = NULL,
body = list(),
raw = FALSE,
asdf = FALSE,
track_total_hits = TRUE,
time_scroll = NULL,
search_path = "_search",
stream_opts = list(),
ignore_unavailable = FALSE,
...
)

```

Arguments

conn	an Elasticsearch connection object, see connect
index	Index name, one or more
type	Document type. Note that type is deprecated in Elasticsearch v7 and greater, and removed in Elasticsearch v8. We will strive to support types for folks using older ES versions
q	The query string (maps to the query_string query, see Query String Query for more details). See https://www.elastic.co/guide/en/elasticsearch/reference/current/query-dsl-query-string-query.html for documentation and examples.
df	(character) The default field to use when no field prefix is defined within the query.
analyzer	(character) The analyzer name to be used when analyzing the query string.
default_operator	(character) The default operator to be used, can be AND or OR. Default: OR
explain	(logical) For each hit, contain an explanation of how scoring of the hits was computed. Default: FALSE
source	(logical) Set to FALSE to disable retrieval of the _source field. You can also retrieve part of the document by using _source_include & _source_exclude (see the body documentation for more details). You can also include a comma-delimited string of fields from the source document that you want back. See also the fields parameter
fields	(character) The selective stored fields of the document to return for each hit. Not specifying any value will cause no fields to return. Note that in Elasticsearch v5 and greater, fields parameter has changed to stored_fields , which is not on by default. You can however, pass fields to source parameter

sort	(character) Sorting to perform. Can either be in the form of fieldName, or fieldName:asc/fieldName:desc. The fieldName can either be an actual field within the document, or the special _score name to indicate sorting based on scores. There can be several sort parameters (order is important).
track_scores	(logical) When sorting, set to TRUE in order to still track scores and return them as part of each hit.
timeout	(numeric) A search timeout, bounding the search request to be executed within the specified time value and bail with the hits accumulated up to that point when expired. Default: no timeout.
terminate_after	(numeric) The maximum number of documents to collect for each shard, upon reaching which the query execution will terminate early. If set, the response will have a boolean field terminated_early to indicate whether the query execution has actually terminated_early. Default: no terminate_after
from	(character) The starting from index of the hits to return. Pass in as a character string to avoid problems with large number conversion to scientific notation. Default: 0
size	(character) The number of hits to return. Pass in as a character string to avoid problems with large number conversion to scientific notation. Default: 10. The default maximum is 10,000 - however, you can change this default maximum by changing the index.max_result_window index level parameter.
search_type	(character) The type of the search operation to perform. Can be query_then_fetch (default) or dfs_query_then_fetch. Types scan and count are deprecated. See Elasticsearch docs for more details on the different types of search that can be performed.
lowercase_expanded_terms	(logical) Should terms be automatically lowercased or not. Default: TRUE.
analyze_wildcard	(logical) Should wildcard and prefix queries be analyzed or not. Default: FALSE.
version	(logical) Print the document version with each document.
lenient	(logical) If TRUE will cause format based failures (like providing text to a numeric field) to be ignored. Default: NULL
body	Query, either a list or json.
raw	(logical) If FALSE (default), data is parsed to list. If TRUE, then raw JSON returned
asdf	(logical) If TRUE, use fromJSON to parse JSON directly to a data.frame. If FALSE (Default), list output is given.
track_total_hits	(logical, numeric) If TRUE will always track the number of hits that match the query accurately. If FALSE will count documents accurately up to 10000 documents. If is.integer will count documents accurately up to the number. Default: TRUE
time_scroll	(character) Specify how long a consistent view of the index should be maintained for scrolled search, e.g., "30s", "1m". See units-time

search_path	(character) The path to use for searching. Default to <code>_search</code> , but in some cases you may already have that in the base url set using <code>connect()</code> , in which case you can set this to <code>NULL</code>
stream_opts	(list) A list of options passed to <code>stream_out</code> - Except that you can't pass <code>x</code> as that's the data that's streamed out, and pass a file path instead of a connection to <code>con</code> . <code>pagesize</code> param doesn't do much as that's more or less controlled by paging with ES.
ignore_unavailable	(logical) What to do if an specified index name doesn't exist. If set to <code>TRUE</code> then those indices are ignored.
...	Curl args passed on to <code>verb-POST</code>

Details

This function name has the "S" capitalized to avoid conflict with the function `base::search`. I hate mixing cases, as I think it confuses users, but in this case it seems necessary.

profile

The Profile API provides detailed timing information about the execution of individual components in a search request. See <https://www.elastic.co/guide/en/elasticsearch/reference/current/search-profile.html> for more information

In a body query, you can set to `profile: true` to enable profiling results. e.g.

```
{
  "profile": true,
  "query" : {
    "match" : { "message" : "some number" }
  }
}
```

References

<https://www.elastic.co/guide/en/elasticsearch/reference/current/search-search.html>
<https://www.elastic.co/guide/en/elasticsearch/reference/current/query-dsl.html>

See Also

[Search_uri\(\)](#) [Search_template\(\)](#) [scroll\(\)](#) [count\(\)](#) [validate\(\)](#) [fielddata\(\)](#)

Examples

```
## Not run:
# make connection object
(x <- connect())

# load some data
if (!index_exists(x, "shakespeare")) {
  shakespeare <- system.file("examples", "shakespeare_data.json",
```

```

    package = "elastic")
  shakespeare <- type_removal(shakespeare)
  invisible(docs_bulk(x, shakespeare))
}
if (!index_exists(x, "gbif")) {
  gbif <- system.file("examples", "gbif_data.json",
    package = "elastic")
  gbif <- type_removal(gbif)
  invisible(docs_bulk(x, gbif))
}
if (!index_exists(x, "plos")) {
  plos <- system.file("examples", "plos_data.json",
    package = "elastic")
  plos <- type_removal(plos)
  invisible(docs_bulk(x, plos))
}

# URI string queries
Search(x, index="shakespeare")
## if you're using an older ES version, you may have types
if (gsub("\\.", "", x$ping()$version$number) < 700) {
  Search(x, index="shakespeare", type="act")
  Search(x, index="shakespeare", type="scene")
  Search(x, index="shakespeare", type="line")
}

## Return certain fields
if (gsub("\\.", "", x$ping()$version$number) < 500) {
  ### ES < v5
  Search(x, index="shakespeare", fields=c('play_name','speaker'))
} else {
  ### ES > v5
  Search(x, index="shakespeare", body = '{
    "_source": ["play_name", "speaker"]
  }')
}

## Search multiple indices
Search(x, index = "gbif")$hits$total$value
Search(x, index = "shakespeare")$hits$total$value
Search(x, index = c("gbif", "shakespeare"))$hits$total$value

## search_type
Search(x, index="shakespeare", search_type = "query_then_fetch")
Search(x, index="shakespeare", search_type = "dfs_query_then_fetch")
### search type "scan" is gone - use time_scroll instead
Search(x, index="shakespeare", time_scroll = "2m")
### search type "count" is gone - use size=0 instead
Search(x, index="shakespeare", size = 0)$hits$total$value

## search exists check
### use size set to 0 and terminate_after set to 1

```

```

### if there are > 0 hits, then there are matching documents
Search(x, index="shakespeare", size = 0, terminate_after = 1)

## sorting
### if ES >5, we need to make sure fielddata is turned on for a field
### before using it for sort
if (gsub("\\.", "", x$ping()$version$number) >= 500) {
  if (index_exists(x, "shakespeare")) index_delete(x, "shakespeare")
  index_create(x, "shakespeare")
  mapping_create(x, "shakespeare", body = '{
    "properties": {
      "speaker": {
        "type": "text",
        "fielddata": true
      }
    }
  }')
}
shakespeare <- system.file("examples", "shakespeare_data.json",
  package = "elastic")
shakespeare <- type_remover(shakespeare)
invisible(docs_bulk(x, shakespeare))
z <- Search(x, index="shakespeare", sort="speaker", size = 30)
vapply(z$hits$hits, function(w) w$_source$speaker, "")
}

if (gsub("\\.", "", x$ping()$version$number) < 500) {
  Search(x, index="shakespeare", type="line", sort="speaker:desc",
    fields='speaker')
  Search(x, index="shakespeare", type="line",
    sort=c("speaker:desc", "play_name:asc"), fields=c('speaker', 'play_name'))
}

## pagination
Search(x, index="shakespeare", size=1)$hits$hits
Search(x, index="shakespeare", size=1, from=1)$hits$hits

## queries
### Search in all fields
Search(x, index="shakespeare", q="york")

### Search in specific fields
Search(x, index="shakespeare", q="speaker:KING HENRY IV")$hits$total$value

### Exact phrase search by wrapping in quotes
Search(x, index="shakespeare", q='speaker:"KING HENRY IV"')$hits$total$value

### can specify operators between multiple words parenthetically
Search(x, index="shakespeare", q="speaker:(HENRY OR ARCHBISHOP)")$hits$total$value

### where the field line_number has no value (or is missing)
Search(x, index="shakespeare", q="_missing_:line_number")$hits$total$value

```

```

### where the field line_number has any non-null value
Search(x, index="shakespeare", q="_exists_:line_number")$hits$total$value

### wildcards, either * or ?
Search(x, index="shakespeare", q="*ay")$hits$total$value
Search(x, index="shakespeare", q="m?y")$hits$total$value

### regular expressions, wrapped in forward slashes
Search(x, index="shakespeare", q="text_entry:/[a-z]/")$hits$total$value

### fuzziness
Search(x, index="shakespeare", q="text_entry:ma~")$hits$total$value
Search(x, index="shakespeare", q="text_entry:the~2")$hits$total$value
Search(x, index="shakespeare", q="text_entry:the~1")$hits$total$value

### Proximity searches
Search(x, index="shakespeare", q='text_entry:"as hath"~5')$hits$total$value
Search(x, index="shakespeare", q='text_entry:"as hath"~10')$hits$total$value

### Ranges, here where line_id value is between 10 and 20
Search(x, index="shakespeare", q="line_id:[10 TO 20]")$hits$total$value

### Grouping
Search(x, index="shakespeare", q="(hath OR as) AND the")$hits$total$value

# Limit number of hits returned with the size parameter
Search(x, index="shakespeare", size=1)

# Give explanation of search in result
Search(x, index="shakespeare", size=1, explain=TRUE)

## terminate query after x documents found
## setting to 1 gives back one document for each shard
Search(x, index="shakespeare", terminate_after=1)
## or set to other number
Search(x, index="shakespeare", terminate_after=2)

## Get version number for each document
Search(x, index="shakespeare", version=TRUE, size=2)

## Get raw data
Search(x, index="shakespeare", raw = TRUE)

## Curl options
### verbose
out <- Search(x, index="shakespeare", verbose = TRUE)

# Query DSL searches - queries sent in the body of the request
## Pass in as an R list

### if ES >5, we need to make sure fielddata is turned on for a field

```

```

### before using it for aggregations
if (gsub("\\.", "", x$ping()$version$number) >= 500) {
  mapping_create(x, "shakespeare", update_all_types = TRUE, body = '{
    "properties": {
      "text_entry": {
        "type": "text",
        "fielddata": true
      }
    }
  }')
  aggs <- list(aggs = list(stats = list(terms = list(field = "text_entry"))))
  Search(x, index="shakespeare", body=aggs)
}

### if ES >5, you don't need to worry about fielddata
if (gsub("\\.", "", x$ping()$version$number) < 500) {
  aggs <- list(aggs = list(stats = list(terms = list(field = "text_entry"))))
  Search(x, index="shakespeare", body=aggs)
}

## or pass in as json query with newlines, easy to read
aggs <- '{
  "aggs": {
    "stats" : {
      "terms" : {
        "field" : "speaker"
      }
    }
  }
}'
Search(x, index="shakespeare", body=aggs, asdf=TRUE, size = 0)

## or pass in collapsed json string
aggs <- '{"aggs":{"stats":{"terms":{"field":"text_entry"}}}}'
Search(x, index="shakespeare", body=aggs)

## Aggregations
### Histograms
aggs <- '{
  "aggs": {
    "latbuckets" : {
      "histogram" : {
        "field" : "decimalLatitude",
        "interval" : 5
      }
    }
  }
}'
Search(x, index="gbif", body=aggs, size=0)

### Histograms w/ more options
aggs <- '{

```

```

    "aggs": {
      "latbuckets" : {
        "histogram" : {
          "field" : "decimalLatitude",
          "interval" : 5,
          "min_doc_count" : 0,
          "extended_bounds" : {
            "min" : -90,
            "max" : 90
          }
        }
      }
    }
  }
}'
Search(x, index="gbif", body=aggs, size=0)

### Ordering the buckets by their doc_count - ascending:
aggs <- '{
  "aggs": {
    "latbuckets" : {
      "histogram" : {
        "field" : "decimalLatitude",
        "interval" : 5,
        "min_doc_count" : 0,
        "extended_bounds" : {
          "min" : -90,
          "max" : 90
        },
        "order" : {
          "_count" : "desc"
        }
      }
    }
  }
}'
out <- Search(x, index="gbif", body=aggs, size=0)
lapply(out$aggregations$latbuckets$buckets, data.frame)

### By default, the buckets are returned as an ordered array. It is also possible to
### request the response as a hash instead keyed by the buckets keys:
aggs <- '{
  "aggs": {
    "latbuckets" : {
      "histogram" : {
        "field" : "decimalLatitude",
        "interval" : 10,
        "keyed" : true
      }
    }
  }
}'
Search(x, index="gbif", body=aggs, size=0)

```

```

# match query
match <- '{"query": {"match" : {"text_entry" : "Two Gentlemen"}}}'
Search(x, index="shakespeare", body=match)

# multi-match (multiple fields that is) query
mmatch <- '{"query": {"multi_match" : {"query" : "henry", "fields": ["text_entry", "play_name"]}}}'
Search(x, index="shakespeare", body=mmatch)

# bool query
mmatch <- '{
  "query": {
    "bool" : {
      "must_not" : {
        "range" : {
          "speech_number" : {
            "from" : 1, "to": 5
          }
        }
      }
    }
  }
}'
Search(x, index="shakespeare", body=mmatch)

# Boosting query
boost <- '{
  "query" : {
    "boosting" : {
      "positive" : {
        "term" : {
          "play_name" : "henry"
        }
      },
      "negative" : {
        "term" : {
          "text_entry" : "thou"
        }
      },
      "negative_boost" : 0.8
    }
  }
}'
Search(x, index="shakespeare", body=boost)

# Fuzzy query
## fuzzy query on numerics
fuzzy <- list(query = list(fuzzy = list(text_entry = "arms")))
Search(x, index="shakespeare", body=fuzzy)$hits$total$value
fuzzy <- list(query = list(fuzzy = list(text_entry = list(value = "arms", fuzziness = 4))))
Search(x, index="shakespeare", body=fuzzy)$hits$total$value

# geoshape query
## not working yet
geo <- list(query = list(geo_shape = list(location = list(shape = list(type = "envelope",
  coordinates = "[[2,10],[10,20]]")))))
geo <- '{
  "query": {
    "geo_shape": {

```

```

      "location": {
        "point": {
          "type": "envelope",
          "coordinates": [[2,0],[2.93,100]]
        }
      }
    }
  }
}'
# Search(x, index="gbifnewgeo", body=geo)

# range query
## with numeric
body <- list(query=list(range=list(decimalLongitude=list(gte=1, lte=3))))
Search(x, 'gbif', body=body)$hits$total$value

body <- list(query=list(range=list(decimalLongitude=list(gte=2.9, lte=10))))
Search(x, 'gbif', body=body)$hits$total$value

## with dates
body <- list(query=list(range=list(eventDate=list(gte="2012-01-01", lte="now"))))
Search(x, 'gbif', body=body)$hits$total$value

body <- list(query=list(range=list(eventDate=list(gte="2014-01-01", lte="now"))))
Search(x, 'gbif', body=body)$hits$total$value

# more like this query (more_like_this can be shortened to mlt)
body <- '{
  "query": {
    "more_like_this": {
      "fields": ["title"],
      "like": "and then",
      "min_term_freq": 1,
      "max_query_terms": 12
    }
  }
}'
Search(x, 'plos', body=body)$hits$total$value

body <- '{
  "query": {
    "more_like_this": {
      "fields": ["abstract","title"],
      "like": "cell",
      "min_term_freq": 1,
      "max_query_terms": 12
    }
  }
}'
Search(x, 'plos', body=body)$hits$total$value

# Highlighting
body <- '{

```

```

    "query": {
      "query_string": {
        "query" : "cell"
      }
    },
    "highlight": {
      "fields": {
        "title": {"number_of_fragments": 2}
      }
    }
  }
}'
out <- Search(x, 'plos', body=body)
out$hits$total$value
sapply(out$hits$hits, function(x) x$_source$title[[1]])

### Common terms query
body <- '{
  "query" : {
    "match": {
      "text_entry": {
        "query": "this is"
      }
    }
  }
}'
Search(x, 'shakespeare', body=body)

## Scrolling search - instead of paging
res <- Search(x, index = 'shakespeare', q="a*", time_scroll="1m")
scroll(x, res$_scroll_id`)

res <- Search(x, index = 'shakespeare', q="a*", time_scroll="5m")
out <- list()
hits <- 1
while(hits != 0){
  res <- scroll(x, res$_scroll_id`)
  hits <- length(res$hits$hits)
  if(hits > 0)
    out <- c(out, res$hits$hits)
}

### Sliced scrolling
#### For scroll queries that return a lot of documents it is possible to
#### split the scroll in multiple slices which can be consumed independently
body1 <- '{
  "slice": {
    "id": 0,
    "max": 2
  },
  "query": {
    "match" : {
      "text_entry" : "a*"
    }
  }
}
```

```

    }
  }'

body2 <- '{
  "slice": {
    "id": 1,
    "max": 2
  },
  "query": {
    "match" : {
      "text_entry" : "a*"
    }
  }
}'

res1 <- Search(x, index = 'shakespeare', time_scroll="1m", body = body1)
res2 <- Search(x, index = 'shakespeare', time_scroll="1m", body = body2)
scroll(x, res1$`_scroll_id`)
scroll(x, res2$`_scroll_id`)

out1 <- list()
hits <- 1
while(hits != 0){
  tmp1 <- scroll(x, res1$`_scroll_id`)
  hits <- length(tmp1$hits$hits)
  if(hits > 0)
    out1 <- c(out1, tmp1$hits$hits)
}

out2 <- list()
hits <- 1
while(hits != 0) {
  tmp2 <- scroll(x, res2$`_scroll_id`)
  hits <- length(tmp2$hits$hits)
  if(hits > 0)
    out2 <- c(out2, tmp2$hits$hits)
}

c(
  lapply(out1, "[[", "_source"),
  lapply(out2, "[[", "_source")
)

# Using filters
## A bool filter
body <- '{
  "query":{
    "bool": {
      "must_not" : {
        "range" : {
          "year" : { "from" : 2011, "to" : 2012 }
        }
      }
    }
  }
}'

```

```

    }
  }
}
}'
Search(x, 'gbif', body = body)$hits$total$value

## Geo filters - fun!
### Note that filters have many geospatial filter options, but queries
### have fewer, and require a geo_shape mapping

body <- '{
  "mappings": {
    "properties": {
      "location" : {"type" : "geo_point"}
    }
  }
}'
index_recreate(x, index='gbifgeopoint', body=body)
path <- system.file("examples", "gbif_geopoint.json",
  package = "elastic")
path <- type_remover(path)
invisible(docs_bulk(x, path))

### Points within a bounding box
body <- '{
  "query":{
    "bool" : {
      "must" : {
        "match_all" : {}
      },
      "filter":{
        "geo_bounding_box" : {
          "location" : {
            "top_left" : {
              "lat" : 60,
              "lon" : 1
            },
            "bottom_right" : {
              "lat" : 40,
              "lon" : 14
            }
          }
        }
      }
    }
  }
}'
out <- Search(x, 'gbifgeopoint', body = body, size = 300)
out$hits$total$value
do.call(rbind, lapply(out$hits$hits, function(x) x$`_source`$location))

### Points within distance of a point

```

```

body <- '{
"query": {
  "bool" : {
    "must" : {
      "match_all" : {}
    },
    "filter" : {
      "geo_distance" : {
        "distance" : "200km",
        "location" : {
          "lon" : 4,
          "lat" : 50
        }
      }
    }
  }
}
}'
out <- Search(x, 'gbifgeopoint', body = body)
out$hits$total$value
do.call(rbind, lapply(out$hits$hits, function(x) x$`_source`$location))

### Points within distance range of a point
body <- '{
"aggs":{
  "points_within_dist" : {
    "geo_distance" : {
      "field": "location",
      "origin" : "4, 50",
      "ranges": [
        {"from" : 200},
        {"to" : 400}
      ]
    }
  }
}
}'
out <- Search(x, 'gbifgeopoint', body = body)
out$hits$total$value
do.call(rbind, lapply(out$hits$hits, function(x) x$`_source`$location))

### Points within a polygon
body <- '{
"query":{
  "bool" : {
    "must" : {
      "match_all" : {}
    },
    "filter":{
      "geo_polygon" : {
        "location" : {
          "points" : [
            [80.0, -20.0], [-80.0, -20.0], [-80.0, 60.0], [40.0, 60.0], [80.0, -20.0]
          ]
        }
      }
    }
  }
}
}'

```

```

    }
  }
}
}'
out <- Search(x, 'gbifgeopoint', body = body)
out$hits$total$value
do.call(rbind, lapply(out$hits$hits, function(x) x`_source`$location))

### Geoshape filters using queries instead of filters
#### Get data with geojson type location data loaded first
body <- '{
  "mappings": {
    "properties": {
      "location" : {"type" : "geo_shape"}
    }
  }
}'
index_recreate(x, index='geoshape', body=body)
path <- system.file("examples", "gbif_geoshape.json",
  package = "elastic")
path <- type_remover(path)
invisible(docs_bulk(x, path))

#### Get data with a square envelope, w/ point defining upper left and the other
#### defining the lower right
body <- '{
  "query":{
    "geo_shape" : {
      "location" : {
        "shape" : {
          "type": "envelope",
          "coordinates": [[-30, 50],[30, 0]]
        }
      }
    }
  }
}'
out <- Search(x, 'geoshape', body = body)
out$hits$total$value

#### Get data with a circle, w/ point defining center, and radius
body <- '{
  "query":{
    "geo_shape" : {
      "location" : {
        "shape" : {
          "type": "circle",
          "coordinates": [-10, 45],
          "radius": "2000km"
        }
      }
    }
  }
}'

```

```

    }
  }'
  out <- Search(x, 'geoshape', body = body)
  out$hits$total$value

#### Use a polygon, w/ point defining center, and radius
body <- '{
  "query":{
    "geo_shape" : {
      "location" : {
        "shape" : {
          "type": "polygon",
          "coordinates": [
            [ [80.0, -20.0], [-80.0, -20.0], [-80.0, 60.0], [40.0, 60.0], [80.0, -20.0] ]
          ]
        }
      }
    }
  }
}'
  out <- Search(x, 'geoshape', body = body)
  out$hits$total$value

# Geofilter with WKT
# format follows "BBOX (minlon, maxlon, maxlat, minlat)"
body <- '{
  "query": {
    "bool" : {
      "must" : {
        "match_all" : {}
      },
      "filter" : {
        "geo_bounding_box" : {
          "location" : {
            "wkt" : "BBOX (1, 14, 60, 40)"
          }
        }
      }
    }
  }
}'
  out <- Search(x, 'gbifgeopoint', body = body)
  out$hits$total$value

# Missing filter
if (gsub("\\.", "", x$ping()$version$number) < 500) {
  ### ES < v5
  body <- '{
    "query":{
      "constant_score" : {

```

```

        "filter" : {
          "missing" : { "field" : "play_name" }
        }
      }
    }
  }'
  Search(x, "shakespeare", body = body)
} else {
  ### ES => v5
  body <- '{
    "query":{
      "bool" : {
        "must_not" : {
          "exists" : {
            "field" : "play_name"
          }
        }
      }
    }
  }'
  Search(x, "shakespeare", body = body)
}

# prefix filter
body <- '{
  "query": {
    "bool": {
      "must": {
        "prefix" : {
          "speaker" : "we"
        }
      }
    }
  }
}'
z <- Search(x, "shakespeare", body = body)
z$hits$total$value
vapply(z$hits$hits, "[", "", c("_source", "speaker"))

# ids filter
if (gsub("\\.", "", x$ping()$version$number) < 500) {
  ### ES < v5
  body <- '{
    "query":{
      "bool": {
        "must": {
          "ids" : {
            "values": ["1","2","10","2000"]
          }
        }
      }
    }
  }'
}

```

```

}'
z <- Search(x, "shakespeare", body = body)
z$hits$total$value
identical(
  c("1","2","10","2000"),
  vapply(z$hits$hits, "[", "", "_id")
)
} else {
body <- '{
  "query":{
    "ids" : {
      "values": ["1","2","10","2000"]
    }
  }
}'
z <- Search(x, "shakespeare", body = body)
z$hits$total$value
identical(
  c("1","2","10","2000"),
  vapply(z$hits$hits, "[", "", "_id")
)
}

# combined prefix and ids filters
if (gsub("\\.", "", x$ping()$version$number) < 500) {
  ### ES < v5
  body <- '{
    "query":{
      "bool" : {
        "should" : {
          "or": [{
            "ids" : {
              "values": ["1","2","3","10","2000"]
            }
          }, {
            "prefix" : {
              "speaker" : "we"
            }
          }
        ]
      }
    }
  }'
  z <- Search(x, "shakespeare", body = body)
  z$hits$total$value
} else {
  ### ES => v5
  body <- '{
    "query":{
      "bool" : {
        "should" : [
          {

```

```

        "ids" : {
          "values": ["1","2","3","10","2000"]
        }
      },
      {
        "prefix" : {
          "speaker" : "we"
        }
      }
    ]
  }
}
}'
z <- Search(x, "shakespeare", body = body)
z$hits$total$value
}

# Suggestions
sugg <- '{
  "query" : {
    "match" : {
      "text_entry" : "late"
    }
  },
  "suggest" : {
    "sugg" : {
      "text" : "late",
      "term" : {
        "field" : "text_entry"
      }
    }
  }
}'
Search(x, index = "shakespeare", body = sugg,
  asdf = TRUE, size = 0)$suggest$sugg$options

# stream data out using jsonlite::stream_out
file <- tempfile()
res <- Search(x, "shakespeare", size = 1000, stream_opts = list(file = file))
head(df <- jsonlite::stream_in(file(file)))
NROW(df)
unlink(file)

# get profile data
body <- '{
  "profile": true,
  "query" : {
    "match" : { "text_entry" : "war" }
  }
}'

```

```

res <- Search(x, "shakespeare", body = body)
res$profile
# time in nanoseconds across each of the shards
vapply(res$profile$shards, function(w) {
  w$searches[[1]]$query[[1]]$time_in_nanos
}, 1)

## End(Not run)

```

searchapis

Overview of search functions

Description

Overview of search functions

Details

Elasticsearch search APIs include the following functions:

- `Search()` - Search using the Query DSL via the body of the request.
- `Search_uri()` - Search using the URI search API only. This may be needed for servers that block POST requests for security, or maybe you don't need complicated requests, in which case URI only requests are suffice.
- `msearch()` - Multi Search - execute several search requests defined in a file passed to msearch
- `search_shards()` - Search shards.
- `count()` - Get counts for various searches.
- `explain()` - Computes a score explanation for a query and a specific document. This can give useful feedback whether a document matches or didn't match a specific query.
- `validate()` - Validate a search
- `field_stats()` - Search field statistics
- `percolate()` - Store queries into an index then, via the percolate API, define documents to retrieve these queries.

More will be added soon.

References

<https://www.elastic.co/guide/en/elasticsearch/reference/current/search.html>

search_shards	<i>Search shards</i>
---------------	----------------------

Description

Search shards

Usage

```
search_shards(
  conn,
  index = NULL,
  raw = FALSE,
  routing = NULL,
  preference = NULL,
  local = NULL,
  ...
)
```

Arguments

conn	an Elasticsearch connection object, see connect()
index	One or more indices
raw	If TRUE (default), data is parsed to list. If FALSE, then raw JSON
routing	A character vector of routing values to take into account when determining which shards a request would be executed against.
preference	Controls a preference of which shard replicas to execute the search request on. By default, the operation is randomized between the shard replicas. See preference for a list of all acceptable values.
local	(logical) Whether to read the cluster state locally in order to determine where shards are allocated instead of using the Master node's cluster state.
...	Curl args passed on to curl::verb-GET

References

<https://www.elastic.co/guide/en/elasticsearch/reference/current/search-shards.html>

Examples

```
## Not run:
# connection setup
(x <- connect())

search_shards(x, index = "plos")
search_shards(x, index = c("plos", "gbif"))
search_shards(x, index = "plos", preference='_primary')
```

```
search_shards(x, index = "plos", preference='_shards:2')

# curl options
search_shards(x, index = "plos", verbose = TRUE)

## End(Not run)
```

Search_template	<i>Search or validate templates</i>
-----------------	-------------------------------------

Description

Search or validate templates

Usage

```
Search_template(conn, body = list(), raw = FALSE, ...)

Search_template_register(conn, template, body = list(), raw = FALSE, ...)

Search_template_get(conn, template, ...)

Search_template_delete(conn, template, ...)

Search_template_render(conn, body = list(), raw = FALSE, ...)
```

Arguments

conn	an Elasticsearch connection object, see connect()
body	Query, either a list or json.
raw	(logical) If FALSE (default), data is parsed to list. If TRUE, then raw JSON returned
...	Curl args passed on to curl::verb-POST
template	(character) a template name

Template search

With Search_template you can search with a template, using mustache templating. Added in Elasticsearch v1.1

Template render

With Search_template_render you validate a template without conducting the search. Added in Elasticsearch v2.0

Pre-registered templates

Register a template with `Search_template_register`. You can get the template with `Search_template_get` and delete the template with `Search_template_delete`

You can also pre-register search templates by storing them in the `config/scripts` directory, in a file using the `.mustache` extension. In order to execute the stored template, reference it by its name under the `template` key, like `"file": "templateName", ...`

References

<https://www.elastic.co/guide/en/elasticsearch/reference/current/search-template.html>

See Also

[Search\(\)](#), [Search_uri\(\)](#)

Examples

```
## Not run:
# connection setup
(x <- connect())

if (!index_exists(x, "iris")) {
  invisible(docs_bulk(x, iris, "iris"))
}

body1 <- '{
  "inline" : {
    "query": { "match" : { "{{my_field}}" : "{{my_value}}" } },
    "size" : "{{my_size}}"
  },
  "params" : {
    "my_field" : "Species",
    "my_value" : "setosa",
    "my_size" : 3
  }
}'
Search_template(x, body = body1)

body2 <- '{
  "inline": {
    "query": {
      "match": {
        "Species": "{{query_string}}"
      }
    }
  },
  "params": {
    "query_string": "versicolor"
  }
}'
```

```

Search_template(x, body = body2)

# pass in a list
mylist <- list(
  inline = list(query = list(match = list(`{{my_field}}` = "{{my_value}}")),
  params = list(my_field = "Species", my_value = "setosa", my_size = 3L)
)
Search_template(x, body = mylist)

## Validating templates w/ Search_template_render()
Search_template_render(x, body = body1)
Search_template_render(x, body = body2)

## pre-registered templates
### register a template
if (x$ses_ver() <= 520) {
  body3 <- '{
    "template": {
      "query": {
        "match": {
          "Species": "{{query_string}}"
        }
      }
    }
  }'
  Search_template_register(x, 'foobar', body = body3)
} else {
  body3 <- '{
    "script": {
      "lang": "mustache",
      "source": {
        "query": {
          "match": {
            "Species": "{{query_string}}"
          }
        }
      }
    }
  }'
  Search_template_register(x, 'foobar', body = body3)
}

### get template
Search_template_get(x, 'foobar')

### use the template
body4 <- '{
  "id": "foobar",
  "params": {
    "query_string": "setosa"
  }
}'
Search_template(x, body = body4)

```

```

### delete the template
Search_template_delete(x, 'foobar')

## End(Not run)

```

Search_uri

Full text search of Elasticsearch with URI search

Description

Full text search of Elasticsearch with URI search

Usage

```

Search_uri(
  conn,
  index = NULL,
  type = NULL,
  q = NULL,
  df = NULL,
  analyzer = NULL,
  default_operator = NULL,
  explain = NULL,
  source = NULL,
  fields = NULL,
  sort = NULL,
  track_scores = NULL,
  timeout = NULL,
  terminate_after = NULL,
  from = NULL,
  size = NULL,
  search_type = NULL,
  lowercase_expanded_terms = NULL,
  analyze_wildcard = NULL,
  version = NULL,
  lenient = NULL,
  raw = FALSE,
  asdf = FALSE,
  track_total_hits = TRUE,
  search_path = "_search",
  stream_opts = list(),
  ignore_unavailable = FALSE,
  ...
)

```

Arguments

conn	an Elasticsearch connection object, see connect
index	Index name, one or more
type	Document type. Note that type is deprecated in Elasticsearch v7 and greater, and removed in Elasticsearch v8. We will strive to support types for folks using older ES versions
q	The query string (maps to the query_string query, see Query String Query for more details). See https://www.elastic.co/guide/en/elasticsearch/reference/current/query-dsl-query-string-query.html for documentation and examples.
df	(character) The default field to use when no field prefix is defined within the query.
analyzer	(character) The analyzer name to be used when analyzing the query string.
default_operator	(character) The default operator to be used, can be AND or OR. Default: OR
explain	(logical) For each hit, contain an explanation of how scoring of the hits was computed. Default: FALSE
source	(logical) Set to FALSE to disable retrieval of the _source field. You can also retrieve part of the document by using _source_include & _source_exclude (see the body documentation for more details). You can also include a comma-delimited string of fields from the source document that you want back. See also the fields parameter
fields	(character) The selective stored fields of the document to return for each hit. Not specifying any value will cause no fields to return. Note that in Elasticsearch v5 and greater, fields parameter has changed to stored_fields , which is not on by default. You can however, pass fields to source parameter
sort	(character) Sorting to perform. Can either be in the form of fieldName, or fieldName:asc/fieldName:desc. The fieldName can either be an actual field within the document, or the special _score name to indicate sorting based on scores. There can be several sort parameters (order is important).
track_scores	(logical) When sorting, set to TRUE in order to still track scores and return them as part of each hit.
timeout	(numeric) A search timeout, bounding the search request to be executed within the specified time value and bail with the hits accumulated up to that point when expired. Default: no timeout.
terminate_after	(numeric) The maximum number of documents to collect for each shard, upon reaching which the query execution will terminate early. If set, the response will have a boolean field terminated_early to indicate whether the query execution has actually terminated_early. Default: no terminate_after
from	(character) The starting from index of the hits to return. Pass in as a character string to avoid problems with large number conversion to scientific notation. Default: 0

size	(character) The number of hits to return. Pass in as a character string to avoid problems with large number conversion to scientific notation. Default: 10. The default maximum is 10,000 - however, you can change this default maximum by changing the <code>index.max_result_window</code> index level parameter.
search_type	(character) The type of the search operation to perform. Can be <code>query_then_fetch</code> (default) or <code>dfs_query_then_fetch</code> . Types <code>scan</code> and <code>count</code> are deprecated. See Elasticsearch docs for more details on the different types of search that can be performed.
lowercase_expanded_terms	(logical) Should terms be automatically lowercased or not. Default: TRUE.
analyze_wildcard	(logical) Should wildcard and prefix queries be analyzed or not. Default: FALSE.
version	(logical) Print the document version with each document.
lenient	(logical) If TRUE will cause format based failures (like providing text to a numeric field) to be ignored. Default: NULL
raw	(logical) If FALSE (default), data is parsed to list. If TRUE, then raw JSON returned
asdf	(logical) If TRUE, use fromJSON to parse JSON directly to a <code>data.frame</code> . If FALSE (Default), list output is given.
track_total_hits	(logical, numeric) If TRUE will always track the number of hits that match the query accurately. If FALSE will count documents accurately up to 10000 documents. If <code>is.integer</code> will count documents accurately up to the number. Default: TRUE
search_path	(character) The path to use for searching. Default to <code>_search</code> , but in some cases you may already have that in the base url set using connect() , in which case you can set this to NULL
stream_opts	(list) A list of options passed to stream_out - Except that you can't pass <code>x</code> as that's the data that's streamed out, and pass a file path instead of a connection to <code>con</code> . <code>pagesize</code> param doesn't do much as that's more or less controlled by paging with ES.
ignore_unavailable	(logical) What to do if an specified index name doesn't exist. If set to TRUE then those indices are ignored.
...	Curl args passed on to verb-POST

See Also

[fielddata\(\)](#)
[Search\(\)](#) [Search_template\(\)](#) [count\(\)](#) [fielddata\(\)](#)

Examples

```
## Not run:
# connection setup
(x <- connect())
```

```

# URI string queries
Search_uri(x, index="shakespeare")
## if you're using an older ES version, you may have types
if (gsub("\\.", "", x$ping()$version$number) < 700) {
  Search_uri(x, index="shakespeare", type="act")
  Search_uri(x, index="shakespeare", type="scene")
  Search_uri(x, index="shakespeare", type="line")
}

## Return certain fields
if (gsub("\\.", "", ping()$version$number) < 500) {
  ### ES < v5
  Search_uri(x, index="shakespeare", fields=c('play_name','speaker'))
} else {
  ### ES > v5
  Search_uri(x, index="shakespeare", source=c('play_name','speaker'))
}

## Search many indices
Search_uri(x, index = "gbif")$hits$total$value
Search_uri(x, index = "shakespeare")$hits$total$value
Search_uri(x, index = c("gbif", "shakespeare"))$hits$total$value

## search_type
## NOTE: If you're in ES V5 or greater, see \code{?felddata}
Search_uri(x, index="shakespeare", search_type = "query_then_fetch")
Search_uri(x, index="shakespeare", search_type = "dfs_query_then_fetch")
# Search_uri(x, index="shakespeare", search_type = "scan") # only when scrolling

## sorting
Search_uri(x, index="shakespeare", sort="text_entry")
if (gsub("\\.", "", x$ping()$version$number) < 500) {
  Search_uri(x, index="shakespeare", sort="speaker:desc", fields='speaker')
  Search_uri(x, index="shakespeare", sort=c("speaker:desc","play_name:asc"),
    fields=c('speaker','play_name'))
}

## pagination
Search_uri(x, index="shakespeare", size=1)$hits$hits
Search_uri(x, index="shakespeare", size=1, from=1)$hits$hits

## queries
### Search in all fields
Search_uri(x, index="shakespeare", q="york")

### Search in specific fields
Search_uri(x, index="shakespeare", q="speaker:KING HENRY IV")$hits$total$value

### Exact phrase search by wrapping in quotes
Search_uri(x, index="shakespeare", q='speaker:"KING HENRY IV"')$hits$total$value

### can specify operators between multiple words parenthetically

```

```

Search_uri(x, index="shakespeare", q="speaker:(HENRY OR ARCHBISHOP)")$hits$total$value

### where the field line_number has no value (or is missing)
Search_uri(x, index="shakespeare", q="_missing_:line_number")$hits$total$value

### where the field line_number has any non-null value
Search_uri(x, index="shakespeare", q="_exists_:line_number")$hits$total$value

### wildcards, either * or ?
Search_uri(x, index="shakespeare", q="*ay")$hits$total$value
Search_uri(x, index="shakespeare", q="m?y")$hits$total$value

### regular expressions, wrapped in forward slashes
Search_uri(x, index="shakespeare", q="text_entry:/[a-z]/")$hits$total$value

### fuzziness
Search_uri(x, index="shakespeare", q="text_entry:ma~")$hits$total$value
Search_uri(x, index="shakespeare", q="text_entry:the~2")$hits$total$value
Search_uri(x, index="shakespeare", q="text_entry:the~1")$hits$total$value

### Proximity searches
Search_uri(x, index="shakespeare", q='text_entry:"as hath"~5')$hits$total$value
Search_uri(x, index="shakespeare", q='text_entry:"as hath"~10')$hits$total$value

### Ranges, here where line_id value is between 10 and 20
Search_uri(x, index="shakespeare", q="line_id:[10 TO 20]")$hits$total$value

### Grouping
Search_uri(x, index="shakespeare", q="(hath OR as) AND the")$hits$total$value

# Limit number of hits returned with the size parameter
Search_uri(x, index="shakespeare", size=1)

# Give explanation of search in result
Search_uri(x, index="shakespeare", size=1, explain=TRUE)

## terminate query after x documents found
## setting to 1 gives back one document for each shard
Search_uri(x, index="shakespeare", terminate_after=1)
## or set to other number
Search_uri(x, index="shakespeare", terminate_after=2)

## Get version number for each document
Search_uri(x, index="shakespeare", version=TRUE, size=2)

## Get raw data
Search_uri(x, index="shakespeare", raw=TRUE)

## Curl options
### verbose
out <- Search_uri(x, index="shakespeare", verbose = TRUE)

## End(Not run)

```

tasks	<i>Elasticsearch tasks endpoints</i>
-------	--------------------------------------

Description

Elasticsearch tasks endpoints

Usage

```
tasks(  
  conn,  
  task_id = NULL,  
  nodes = NULL,  
  actions = NULL,  
  parent_task_id = NULL,  
  detailed = FALSE,  
  group_by = NULL,  
  wait_for_completion = FALSE,  
  timeout = NULL,  
  raw = FALSE,  
  ...  
)  
  
tasks_cancel(  
  conn,  
  node_id = NULL,  
  task_id = NULL,  
  nodes = NULL,  
  actions = NULL,  
  parent_task_id = NULL,  
  detailed = FALSE,  
  group_by = NULL,  
  wait_for_completion = FALSE,  
  timeout = NULL,  
  raw = FALSE,  
  ...  
)
```

Arguments

- conn an Elasticsearch connection object, see [connect\(\)](#)
- task_id a task id
- nodes (character) The nodes
- actions (character) Actions
- parent_task_id (character) A parent task ID

detailed (character) get detailed results. Default: FALSE
group_by (character) "nodes" (default, i.e., NULL) or "parents"
wait_for_completion (logical) wait for completion. Default: FALSE
timeout (integer) timeout time
raw If TRUE (default), data is parsed to list. If FALSE, then raw JSON.
... Curl args passed on to [crul::verb-GET](#) or [crul::verb-POST](#)
node_id a node id

References

<https://www.elastic.co/guide/en/elasticsearch/reference/current/tasks.html>

Examples

```
## Not run:  
x <- connect()  
  
tasks(x)  
# tasks(x, parent_task_id = "1234")  
  
# delete a task  
# tasks_cancel(x)  
  
## End(Not run)
```

termvectors	<i>Termvectors</i>
-------------	--------------------

Description

Termvectors

Usage

```
termvectors(  
  conn,  
  index,  
  type = NULL,  
  id = NULL,  
  body = list(),  
  pretty = TRUE,  
  field_statistics = TRUE,  
  fields = NULL,  
  offsets = TRUE,  
  parent = NULL,  
  payloads = TRUE,
```

```

    positions = TRUE,
    realtime = TRUE,
    preference = "random",
    routing = NULL,
    term_statistics = FALSE,
    version = NULL,
    version_type = NULL,
    ...
)

```

Arguments

<code>conn</code>	an Elasticsearch connection object, see connect()
<code>index</code>	(character) The index in which the document resides.
<code>type</code>	(character) The type of the document. optional
<code>id</code>	(character) The id of the document, when not specified a doc param should be supplied.
<code>body</code>	(character) Define parameters and or supply a document to get termvectors for
<code>pretty</code>	(logical) pretty print. Default: TRUE
<code>field_statistics</code>	(character) Specifies if document count, sum of document frequencies and sum of total term frequencies should be returned. Default: TRUE
<code>fields</code>	(character) A comma-separated list of fields to return.
<code>offsets</code>	(character) Specifies if term offsets should be returned. Default: TRUE
<code>parent</code>	(character) Parent id of documents.
<code>payloads</code>	(character) Specifies if term payloads should be returned. Default: TRUE
<code>positions</code>	(character) Specifies if term positions should be returned. Default: TRUE
<code>realtime</code>	(character) Specifies if request is real-time as opposed to near-real-time (Default: TRUE).
<code>preference</code>	(character) Specify the node or shard the operation should be performed on (Default: random).
<code>routing</code>	(character) Specific routing value.
<code>term_statistics</code>	(character) Specifies if total term frequency and document frequency should be returned. Default: FALSE
<code>version</code>	(character) Explicit version number for concurrency control
<code>version_type</code>	(character) Specific version type, valid choices are: 'internal', 'external', 'external_gte', 'force'
<code>...</code>	Curl args passed on to curl::verb-POST

Details

Returns information and statistics on terms in the fields of a particular document. The document could be stored in the index or artificially provided by the user (Added in 1.4). Note that for documents stored in the index, this is a near realtime API as the term vectors are not available until the next refresh.

References

<https://www.elastic.co/guide/en/elasticsearch/reference/current/docs-termvectors.html>

See Also

[mtermvectors\(\)](#)

Examples

```
## Not run:
x <- connect()

if (!index_exists(x, 'plos')) {
  plosdat <- system.file("examples", "plos_data.json",
    package = "elastic")
  plosdat <- type_remover(plosdat)
  invisible(docs_bulk(x, plosdat))
}
if (!index_exists(x, 'omdb')) {
  omdb <- system.file("examples", "omdb.json", package = "elastic")
  omdb <- type_remover(omdb)
  invisible(docs_bulk(x, omdb))
}

body <- '{
  "fields" : ["title"],
  "offsets" : true,
  "positions" : true,
  "term_statistics" : true,
  "field_statistics" : true
}'
termvectors(x, 'plos', id = 29, body = body)

body <- '{
  "fields" : ["Plot"],
  "offsets" : true,
  "positions" : true,
  "term_statistics" : true,
  "field_statistics" : true
}'
termvectors(x, 'omdb', id = Search(x, "omdb", size=1)$hits$hits[[1]]$`_id`,
  body = body)

## End(Not run)
```

Description

Tokenizer operations

Usage

```
tokenizer_set(conn, index, body, ...)
```

Arguments

conn	an Elasticsearch connection object, see connect()
index	(character) A character vector of index names
body	Query, either a list or json.
...	Curl options passed on to curl::HttpClient

Author(s)

Scott Chamberlain myrmecocystus@gmail.com

References

<https://www.elastic.co/guide/en/elasticsearch/reference/current/analysis-tokenizers.html>

Examples

```
## Not run:
# connection setup
(x <- connect())

# set tokenizer

## NGram tokenizer
body <- '{
  "settings" : {
    "analysis" : {
      "analyzer" : {
        "my_ngram_analyzer" : {
          "tokenizer" : "my_ngram_tokenizer"
        }
      },
      "tokenizer" : {
        "my_ngram_tokenizer" : {
          "type" : "nGram",
          "min_gram" : "2",
          "max_gram" : "3",
          "token_chars": [ "letter", "digit" ]
        }
      }
    }
  }
}'
```

```
if (index_exists('test1')) index_delete('test1')
tokenizer_set(index = "test1", body=body)
index_analyze(text = "hello world", index = "test1",
  analyzer='my_ngram_analyzer')

## End(Not run)
```

type_remover*Utility function to remove 'type' from bulk load files*

Description

Types are being removed from Elasticsearch. This little function aims to help remove "_type" fields from bulk newline-delimited JSON files. See Details.

Usage

```
type_remover(file)
```

Arguments

file (character) a file path, required

Details

Looks for any lines that have an "index" key, then drops any "_type" keys in the hash given by the "index" key.

You can of course manually modify these files as an alternative, in a text editor or with command line tools like sed, etc.

Value

a file path for a temporary file with the types removed

Examples

```
## Not run:
z <- system.file("examples/omdb.json", package = "elastic")
readLines(z, 6)
ff <- type_remover(z)
readLines(ff, 6)
unlink(ff)

## End(Not run)
```

units-distance	<i>Distance units</i>
----------------	-----------------------

Description

Wherever distances need to be specified, such as the distance parameter in the Geo Distance Filter), the default unit if none is specified is the meter. Distances can be specified in other units, such as "1km" or "2mi" (2 miles).

Details

mi or miles	Mile
yd or yards	Yard
ft or feet	Feet
in or inch	Inch
km or kilometers	Kilometer
m or meters	Meter
cm or centimeters	Centimeter
mm or millimeters	Millimeter
NM, nmi or nauticalmiles	Nautical mile

The precision parameter in the Geohash Cell Filter accepts distances with the above units, but if no unit is specified, then the precision is interpreted as the length of the geohash.

See Also

[units-time](#)

units-time	<i>Time units</i>
------------	-------------------

Description

Whenever durations need to be specified, eg for a timeout parameter, the duration can be specified as a whole number representing time in milliseconds, or as a time value like 2d for 2 days. The supported units are:

Details

y	Year
M	Month
w	Week
d	Day
h	Hour
m	Minute
s	Second

See Also

[units-distance](#)

validate	<i>Validate a search</i>
----------	--------------------------

Description

Validate a search

Usage

```
validate(conn, index, type = NULL, ...)
```

Arguments

conn	an Elasticsearch connection object, see connect()
index	Index name. Required.
type	Document type. Optional.
...	Additional args passed on to Search()

See Also

[Search\(\)](#)

Examples

```
## Not run:
x <- connect()

if (!index_exists(x, "twitter")) index_create(x, "twitter")
docs_create(x, 'twitter', id=1, body = list(
  "user" = "foobar",
  "post_date" = "2014-01-03",
  "message" = "trying out Elasticsearch"
)
)
validate(x, "twitter", q='user:foobar')
validate(x, "twitter", q='user:foobar')

body <- '{
  "query" : {
    "bool" : {
      "must" : {
        "query_string" : {
          "query" : "*:*"
        }
      },
      "filter" : {
        "term" : { "user" : "kimchy" }
      }
    }
  }
}'
validate(x, "twitter", body = body)

## End(Not run)
```

Index

* bulk-functions

- docs_bulk, 18
- docs_bulk_create, 23
- docs_bulk_delete, 25
- docs_bulk_index, 26
- docs_bulk_prep, 28
- docs_bulk_update, 32

alias, 3

- alias_create(alias), 3
- alias_delete(alias), 3
- alias_exists(alias), 3
- alias_get(alias), 3
- alias_rename(alias), 3
- aliases_get(alias), 3

base::cat(), 9, 78

cat, 5

- cat_(cat), 5
- cat_(), 9
- cat_aliases(cat), 5
- cat_allocation(cat), 5
- cat_count(cat), 5
- cat_fielddata(cat), 5
- cat_health(cat), 5
- cat_indices(cat), 5
- cat_master(cat), 5
- cat_nodeattrs(cat), 5
- cat_nodes(cat), 5
- cat_pending_tasks(cat), 5
- cat_plugins(cat), 5
- cat_recovery(cat), 5
- cat_segments(cat), 5
- cat_shards(cat), 5
- cat_thread_pool(cat), 5

cluster, 10

- cluster_health(cluster), 11
- cluster_pending_tasks(cluster), 11
- cluster_reroute(cluster), 11

- cluster_settings(cluster), 11
- cluster_state(cluster), 11
- cluster_stats(cluster), 11
- connect, 14, 93, 118
- connect(), 3, 9, 12, 17, 18, 23, 25, 27, 32, 34, 36, 38, 40, 42, 44, 46, 50, 52, 54, 55, 57, 61, 67, 70, 73, 74, 78, 80, 85, 87, 88, 95, 113, 114, 119, 122, 124, 126, 129
- count, 16
- count(), 95, 112, 119
- curl::HttpClient, 9, 19, 24, 26, 27, 33, 35, 36, 41, 45, 52, 55, 61, 126
- curl::verb-DELETE, 3, 67
- curl::verb-GET, 3, 17, 54, 67, 70, 78, 85, 113, 123
- curl::verb-HEAD, 3, 70
- curl::verb-POST, 3, 12, 38, 47, 67, 73, 75, 87, 88, 114, 123, 124
- curl::verb-PUT, 67, 70

docs_bulk, 18, 24, 26, 28, 29, 33

- docs_bulk(), 24, 26, 28, 33, 48, 73
- docs_bulk_create, 20, 23, 26, 28, 29, 33
- docs_bulk_delete, 20, 24, 25, 28, 29, 33
- docs_bulk_index, 20, 24, 26, 26, 29, 33
- docs_bulk_prep, 20, 24, 26, 28, 28, 33
- docs_bulk_update, 20, 24, 26, 28, 29, 32
- docs_create, 34
- docs_create(), 48
- docs_delete, 36
- docs_delete(), 48
- docs_delete_by_query, 37
- docs_delete_by_query(), 47
- docs_get, 40
- docs_get(), 48
- docs_mget, 42
- docs_mget(), 48
- docs_update, 44
- docs_update_by_query, 46

- docs_update_by_query(), 38
- documents, 48
- elastic, 49
- elastic-defunct, 51
- elastic-package (elastic), 49
- explain, 51
- explain(), 112
- field_caps, 54
- field_caps(), 56
- field_mapping_get (mapping), 69
- field_stats, 55
- field_stats(), 54, 112
- fielddata, 53
- fielddata(), 95, 119
- fromJSON, 55, 94, 119
- HttpClient, 42
- index_analyze (indices), 58
- index_clear_cache (indices), 58
- index_close (indices), 58
- index_create (indices), 58
- index_create(), 62
- index_delete (indices), 58
- index_exists (indices), 58
- index_flush (indices), 58
- index_forcemerge (indices), 58
- index_get (indices), 58
- index_open (indices), 58
- index_optimize (indices), 58
- index_recovery (indices), 58
- index_recovery(), 62
- index_recreate (indices), 58
- index_segments (indices), 58
- index_settings (indices), 58
- index_settings_update (indices), 58
- index_shrink (indices), 58
- index_stats (indices), 58
- index_status(), 51
- index_template, 56
- index_template_delete (index_template), 56
- index_template_exists (index_template), 56
- index_template_get (index_template), 56
- index_template_put (index_template), 56
- index_upgrade (indices), 58
- indices, 58
- ingest, 66
- invisible(), 20
- jsonlite::fromJSON(), 73, 88
- jsonlite::stream_out(), 88
- jsonlite::toJSON(), 19, 27, 29, 33
- mapping, 69
- mapping_create (mapping), 69
- mapping_delete(), 51
- mapping_get (mapping), 69
- mlt(), 51
- msearch, 73
- msearch(), 112
- mtermvectors, 74
- mtermvectors(), 125
- nodes, 77
- nodes_hot_threads (nodes), 77
- nodes_hot_threads(), 78
- nodes_info (nodes), 77
- nodes_shutdown(), 51
- nodes_stats (nodes), 77
- percolate, 79
- percolate(), 112
- percolate_count (percolate), 79
- percolate_delete (percolate), 79
- percolate_list (percolate), 79
- percolate_match (percolate), 79
- percolate_register (percolate), 79
- ping, 85
- pipeline_attachment (ingest), 66
- pipeline_create (ingest), 66
- pipeline_delete (ingest), 66
- pipeline_get (ingest), 66
- pipeline_simulate (ingest), 66
- preference, 86, 113
- reindex, 86
- scroll, 88
- scroll(), 95
- scroll_clear (scroll), 88
- Search, 92
- Search(), 17, 50, 51, 57, 73, 80, 88, 89, 112, 115, 119, 129
- search_shards, 113
- search_shards(), 112

Search_template, [114](#)
Search_template(), [95](#), [119](#)
Search_template_delete
 (Search_template), [114](#)
Search_template_get(Search_template),
 [114](#)
Search_template_register
 (Search_template), [114](#)
Search_template_render
 (Search_template), [114](#)
Search_uri, [117](#)
Search_uri(), [17](#), [50](#), [73](#), [95](#), [112](#), [115](#)
searchapis, [112](#)
stream_out, [95](#), [119](#)
suppressWarnings(), [15](#)

tasks, [122](#)
tasks_cancel(tasks), [122](#)
termvectors, [123](#)
termvectors(), [75](#)
tokenizer_set, [125](#)
type_exists(mapping), [69](#)
type_removal, [127](#)
type_removal(), [20](#)

units-distance, [128](#), [129](#)
units-time, [88](#), [94](#), [128](#), [128](#)

validate, [129](#)
validate(), [95](#), [112](#)