
Stable Baselines Documentation

Release 2.2.0

Stable Baselines Contributors

Dec 10, 2018

1	Main differences with OpenAI Baselines	3
1.1	Installation	3
1.2	Getting Started	5
1.3	RL Algorithms	6
1.4	Examples	6
1.5	Vectorized Environments	13
1.6	Using Custom Environments	17
1.7	Custom Policy Network	18
1.8	Tensorboard Integration	20
1.9	Base RL Class	22
1.10	Policy Networks	24
1.11	A2C	30
1.12	ACER	33
1.13	ACKTR	36
1.14	DDPG	39
1.15	DQN	50
1.16	GAIL	58
1.17	HER	61
1.18	PPO1	63
1.19	PPO2	66
1.20	TRPO	69
1.21	Probability Distributions	71
1.22	Tensorflow Utils	79
1.23	Command Utils	82
1.24	Schedules	84
1.25	Changelog	85
1.26	Plotting Results	89
2	Citing Stable Baselines	91
3	Contributing	93
4	Indices and tables	95
	Python Module Index	97

Stable Baselines is a set of improved implementations of Reinforcement Learning (RL) algorithms based on OpenAI Baselines.

Github repository: <https://github.com/hill-a/stable-baselines>

You can read a detailed presentation of Stable Baselines in the Medium article: [link](#)

Main differences with OpenAI Baselines

This toolset is a fork of OpenAI Baselines, with a major structural refactoring, and code cleanups:

- Unified structure for all algorithms
- PEP8 compliant (unified code style)
- Documented functions and classes
- More tests & more code coverage

1.1 Installation

1.1.1 Prerequisites

Baselines requires python3 (≥ 3.5) with the development headers. You'll also need system packages CMake, OpenMPI and zlib. Those can be installed as follows

Ubuntu

```
sudo apt-get update && sudo apt-get install cmake libopenmpi-dev python3-dev zlib1g-  
↳dev
```

Mac OS X

Installation of system packages on Mac requires [Homebrew](#). With Homebrew installed, run the following:

```
brew install cmake openmpi
```

1.1.2 Stable Release

```
pip install stable-baselines
```

1.1.3 Bleeding-edge version

With support for running tests and building the documentation.

```
git clone https://github.com/hill-a/stable-baselines && cd stable-baselines
pip install -e .[docs,tests]
```

1.1.4 Using Docker Images

Use Built Images

GPU image (requires [nvidia-docker](#)):

```
docker pull araffin/stable-baselines
```

CPU only:

```
docker pull araffin/stable-baselines-cpu
```

Build the Docker Images

Build GPU image (with [nvidia-docker](#)):

```
docker build . -f docker/Dockerfile.gpu -t stable-baselines
```

Build CPU image:

```
docker build . -f docker/Dockerfile.cpu -t stable-baselines-cpu
```

Note: if you are using a proxy, you need to pass extra params during build and do some [tweaks](#):

```
--network=host --build-arg HTTP_PROXY=http://your.proxy.fr:8080/ --build-arg http_
↪ proxy=http://your.proxy.fr:8080/ --build-arg HTTPS_PROXY=https://your.proxy.fr:8080/
↪ --build-arg https_proxy=https://your.proxy.fr:8080/
```

Run the images (CPU/GPU)

Run the [nvidia-docker](#) GPU image

```
docker run -it --runtime=nvidia --rm --network host --ipc=host --name test --mount_
↪ src="$ (pwd) ",target=/root/code/stable-baselines,type=bind araffin/stable-baselines_
↪ bash -c 'cd /root/code/stable-baselines/ && pytest tests/'
```

Or, with the shell file:

```
./run_docker_gpu.sh pytest tests/
```

Run the docker CPU image

```
docker run -it --rm --network host --ipc=host --name test --mount src="$(pwd)",
↳target=/root/code/stable-baselines,type=bind araffin/stable-baselines-cpu bash -c
↳'cd /root/code/stable-baselines/ && pytest tests/'
```

Or, with the shell file:

```
./run_docker_cpu.sh pytest tests/
```

Explanation of the docker command:

- `docker run -it` create an instance of an image (=container), and run it interactively (so `ctrl+c` will work)
- `--rm` option means to remove the container once it exits/stops (otherwise, you will have to use `docker rm`)
- `--network host` don't use network isolation, this allow to use tensorboard/visdom on host machine
- `--ipc=host` Use the host system's IPC namespace. IPC (POSIX/SysV IPC) namespace provides separation of named shared memory segments, semaphores and message queues.
- `--name test` give explicitly the name `test` to the container, otherwise it will be assigned a random name
- `--mount src=...` give access of the local directory (`pwd` command) to the container (it will be map to `/root/code/stable-baselines`), so all the logs created in the container in this folder will be kept
- `bash -c '...'` Run command inside the docker image, here run the tests (`pytest tests/`)

1.2 Getting Started

Most of the library tries to follow a sklearn-like syntax for the Reinforcement Learning algorithms.

Here is a quick example of how to train and run PPO2 on a cartpole environment:

```
import gym

from stable_baselines.common.policies import MlpPolicy
from stable_baselines.common.vec_env import DummyVecEnv
from stable_baselines import PPO2

env = gym.make('CartPole-v1')
env = DummyVecEnv([lambda: env]) # The algorithms require a vectorized environment,
↳to run

model = PPO2(MlpPolicy, env, verbose=1)
model.learn(total_timesteps=10000)

obs = env.reset()
for i in range(1000):
    action, _states = model.predict(obs)
    obs, rewards, dones, info = env.step(action)
    env.render()
```

Or just train a model with a one liner if the environment is registered in Gym and if the policy is registered:

```
from stable_baselines import PPO2

model = PPO2('MlpPolicy', 'CartPole-v1').learn(10000)
```

Fig. 1: Define and train a RL agent in one line of code!

1.3 RL Algorithms

This table displays the rl algorithms that are implemented in the stable baselines project, along with some useful characteristics: support for recurrent policies, discrete/continuous actions, multiprocessing.

Name	Refactored ¹	Recurrent	Box	Discrete	Multi Processing
A2C					
ACER			⁵		
ACKTR			⁵		
DDPG					
DQN					
GAIL ²					⁴
PPO1					⁴
PPO2					
TRPO					⁴

Actions `gym.spaces`:

- `Box`: A N-dimensional box that contains every point in the action space.
- `Discrete`: A list of possible actions, where each timestep only one of the actions can be used.
- `MultiDiscrete`: A list of possible actions, where each timestep only one action of each discrete set can be used.
- `MultiBinary`: A list of possible actions, where each timestep any of the actions can be used in any combination.

1.4 Examples

1.4.1 Try it online with Colab Notebooks!

All the following examples can be executed online using Google colab notebooks:

- [Getting Started](#)
- [Training, Saving, Loading](#)
- [Multiprocessing](#)
- [Monitor Training and Plotting](#)
- [Atari Games](#)
- [Breakout](#) (trained agent included)

¹ Whether or not the algorithm has been refactored to fit the `BaseRLModel` class.

⁵ TODO, in project scope.

² Only implemented for TRPO.

⁴ Multi Processing with [MPI](#).

1.4.2 Basic Usage: Training, Saving, Loading

In the following example, we will train, save and load an A2C model on the Lunar Lander environment.

Try it in a  notebook

Fig. 2: Lunar Lander Environment

Note: LunarLander requires the python package *box2d*. You can install it using `apt install swig` and then `pip install box2d box2d-kengz`

```
import gym

from stable_baselines.common.policies import MlpPolicy
from stable_baselines.common.vec_env import DummyVecEnv
from stable_baselines import A2C

# Create and wrap the environment
env = gym.make('LunarLander-v2')
env = DummyVecEnv([lambda: env])

model = A2C(MlpPolicy, env, ent_coef=0.1, verbose=1)
# Train the agent
model.learn(total_timesteps=100000)
# Save the agent
model.save("a2c_lunar")
del model # delete trained model to demonstrate loading

# Load the trained agent
model = A2C.load("a2c_lunar")

# Enjoy trained agent
obs = env.reset()
for i in range(1000):
    action, _states = model.predict(obs)
    obs, rewards, dones, info = env.step(action)
    env.render()
```

1.4.3 Multiprocessing: Unleashing the Power of Vectorized Environments

Try it in a  notebook

Fig. 3: CartPole Environment

```
import gym
import numpy as np

from stable_baselines.common.policies import MlpPolicy
from stable_baselines.common.vec_env import SubprocVecEnv
from stable_baselines.common import set_global_seeds
from stable_baselines import ACKTR

def make_env(env_id, rank, seed=0):
    """
    Utility function for multiprocessed env.

    :param env_id: (str) the environment ID
    :param num_env: (int) the number of environments you wish to have in subprocesses
    :param seed: (int) the initial seed for RNG
    :param rank: (int) index of the subprocess
    """
    def _init():
        env = gym.make(env_id)
        env.seed(seed + rank)
        return env
    set_global_seeds(seed)
    return _init

env_id = "CartPole-v1"
num_cpu = 4 # Number of processes to use
# Create the vectorized environment
env = SubprocVecEnv([make_env(env_id, i) for i in range(num_cpu)])

model = ACKTR(MlpPolicy, env, verbose=1)
model.learn(total_timesteps=25000)

obs = env.reset()
for _ in range(1000):
    action, _states = model.predict(obs)
    obs, rewards, dones, info = env.step(action)
    env.render()
```

1.4.4 Using Callback: Monitoring Training

You can define a custom callback function that will be called inside the agent. This could be useful when you want to monitor training, for instance display live learning curves in Tensorboard (or in Visdom) or save the best agent.

Try it in a  notebook

```
import os

import gym
import numpy as np
import matplotlib.pyplot as plt

from stable_baselines.ddpg.policies import MlpPolicy
from stable_baselines.common.vec_env.dummy_vec_env import DummyVecEnv
from stable_baselines.bench import Monitor
```

(continues on next page)

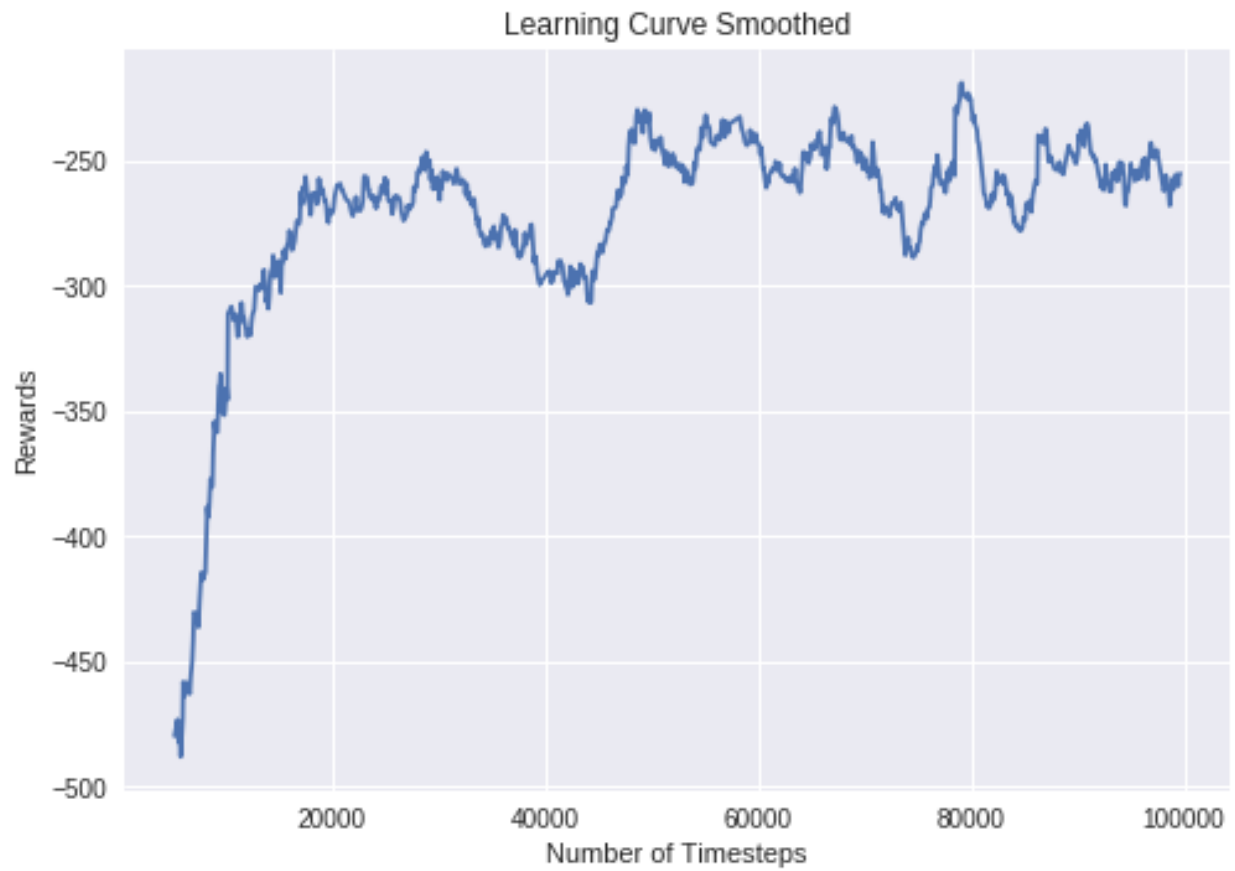


Fig. 4: Learning curve of DDPG on LunarLanderContinuous environment

(continued from previous page)

```

from stable_baselines.results_plotter import load_results, ts2xy
from stable_baselines import DDPG
from stable_baselines.ddpg.noise import AdaptiveParamNoiseSpec

best_mean_reward, n_steps = -np.inf, 0

def callback(_locals, _globals):
    """
    Callback called at each step (for DQN and others) or after n steps (see ACER or PPO2)
    :param _locals: (dict)
    :param _globals: (dict)
    """
    global n_steps, best_mean_reward
    # Print stats every 1000 calls
    if (n_steps + 1) % 1000 == 0:
        # Evaluate policy performance
        x, y = ts2xy(load_results(log_dir), 'timesteps')
        if len(x) > 0:
            mean_reward = np.mean(y[-100:])
            print(x[-1], 'timesteps')
            print("Best mean reward: {:.2f} - Last mean reward per episode: {:.2f}".
                  ↪format(best_mean_reward, mean_reward))

            # New best model, you could save the agent here
            if mean_reward > best_mean_reward:
                best_mean_reward = mean_reward
                # Example for saving best model
                print("Saving new best model")
                _locals['self'].save(log_dir + 'best_model.pkl')

    n_steps += 1
    return False

# Create log dir
log_dir = "/tmp/gym/"
os.makedirs(log_dir, exist_ok=True)

# Create and wrap the environment
env = gym.make('LunarLanderContinuous-v2')
env = Monitor(env, log_dir, allow_early_resets=True)
env = DummyVecEnv([lambda: env])

# Add some param noise for exploration
param_noise = AdaptiveParamNoiseSpec(initial_stddev=0.2, desired_action_stddev=0.2)
model = DDPG(MlpPolicy, env, param_noise=param_noise, memory_limit=int(1e6), ↪
            ↪verbose=0)
# Train the agent
model.learn(total_timesteps=200000, callback=callback)

```

1.4.5 Atari Games

Fig. 5: Trained A2C agent on Breakout

Fig. 6: Pong Environment

Training a RL agent on Atari games is straightforward thanks to `make_atari_env` helper function. It will do all

the preprocessing and multiprocessing for you.

Try it in a  notebook

```
from stable_baselines.common.cmd_util import make_atari_env
from stable_baselines.common.policies import CnnPolicy
from stable_baselines.common.vec_env import VecFrameStack
from stable_baselines import ACER

# There already exists an environment generator
# that will make and wrap atari environments correctly.
# Here we are also multiprocessing training (num_env=4 => 4 processes)
env = make_atari_env('PongNoFrameskip-v4', num_env=4, seed=0)
# Frame-stacking with 4 frames
env = VecFrameStack(env, n_stack=4)

model = ACER(CnnPolicy, env, verbose=1)
model.learn(total_timesteps=25000)

obs = env.reset()
while True:
    action, _states = model.predict(obs)
    obs, rewards, dones, info = env.step(action)
    env.render()
```

1.4.6 Mujoco: Normalizing input features

Normalizing input features may be essential to successful training of an RL agent (by default, images are scaled but not other types of input), for instance when training on [Mujoco](#). For that, a wrapper exists and will compute a running average and standard deviation of input features (it can do the same for rewards).

Note: We cannot provide a notebook for this example because Mujoco is a proprietary engine and requires a license.

```
import gym

from stable_baselines.common.policies import MlpPolicy
from stable_baselines.common.vec_env import DummyVecEnv, VecNormalize
from stable_baselines import PPO2

env = DummyVecEnv([lambda: gym.make("Reacher-v2")])
# Automatically normalize the input features
env = VecNormalize(env, norm_obs=True, norm_reward=False,
                  clip_obs=10.)

model = PPO2(MlpPolicy, env)
model.learn(total_timesteps=2000)

# Don't forget to save the running average when saving the agent
```

(continues on next page)

(continued from previous page)

```
log_dir = "/tmp/"
model.save(log_dir + "ppo_reacher")
env.save_running_average(log_dir)
```

1.4.7 Custom Policy Network

Stable baselines provides default policy networks for images (CNNPolicies) and other type of inputs (MlpPolicies). However, you can also easily define a custom architecture for the policy network (see [custom policy section](#)):

```
import gym

from stable_baselines.common.policies import FeedForwardPolicy
from stable_baselines.common.vec_env import DummyVecEnv
from stable_baselines import A2C

# Custom MLP policy of three layers of size 128 each
class CustomPolicy(FeedForwardPolicy):
    def __init__(self, *args, **kwargs):
        super(CustomPolicy, self).__init__(*args, **kwargs,
                                           layers=[128, 128, 128],
                                           feature_extraction="mlp")

# Create and wrap the environment
env = gym.make('LunarLander-v2')
env = DummyVecEnv([lambda: env])

model = A2C(CustomPolicy, env, verbose=1)
# Train the agent
model.learn(total_timesteps=100000)
```

1.4.8 Continual Learning

You can also move from learning on one environment to another for [continual learning](#) (PPO2 on DemonAttack-v0, then transferred on SpaceInvaders-v0):

```
from stable_baselines.common.cmd_util import make_atari_env
from stable_baselines.common.policies import CnnPolicy
from stable_baselines import PPO2

# There already exists an environment generator
# that will make and wrap atari environments correctly
env = make_atari_env('DemonAttackNoFrameskip-v4', num_env=8, seed=0)

model = PPO2(CnnPolicy, env, verbose=1)
model.learn(total_timesteps=10000)

obs = env.reset()
for i in range(1000):
    action, _states = model.predict(obs)
    obs, rewards, dones, info = env.step(action)
    env.render()

# The number of environments must be identical when changing environments
```

(continues on next page)

(continued from previous page)

```

env = make_atari_env('SpaceInvadersNoFrameskip-v4', num_env=8, seed=0)

# change env
model.set_env(env)
model.learn(total_timesteps=10000)

obs = env.reset()
while True:
    action, _states = model.predict(obs)
    obs, rewards, dones, info = env.step(action)
    env.render()

```

1.4.9 Bonus: Make a GIF of a Trained Agent

Note: For Atari games, you need to use a screen recorder such as [Kazam](#). And then convert the video using [ffmpeg](#)

```

import imageio
import numpy as np

from stable_baselines.common.policies import MlpPolicy
from stable_baselines import A2C

model = A2C(MlpPolicy, "LunarLander-v2").learn(100000)

images = []
obs = model.env.reset()
img = model.env.render(mode='rgb_array')
for i in range(350):
    images.append(img)
    action, _ = model.predict(obs)
    obs, _, _, _ = model.env.step(action)
    img = model.env.render(mode='rgb_array')

imageio.mimsave('lander_a2c.gif', [np.array(img[0]) for i, img in enumerate(images)
↪ if i%2 == 0], fps=29)

```

1.5 Vectorized Environments

Vectorized Environments are a way to multiprocessing training. Instead of training a RL agent on 1 environment, it allows to train it on n environments using n processes. Because of that, *actions* passed to the environment are now a vector (of dimension n). It is the same for *observations*, *rewards* and end of episode signals (*dones*).

Note: Vectorized environments are required when using wrappers for frame-stacking or normalization.

Note: When using vectorized environments, the environments are automatically resetted at the end of each episode.

Warning: It seems that Windows users are experiencing issues with SubprocVecEnv. We recommend to use the docker image in that case. (See [Issue #42](#))

1.5.1 DummyVecEnv

class `stable_baselines.common.vec_env.DummyVecEnv` (*env_fns*)

Creates a simple vectorized wrapper for multiple environments

Parameters `env_fns` – ([Gym Environment]) the list of environments to vectorize

close ()

Clean up the environment's resources.

env_method (*method_name*, **method_args*, ***method_kwargs*)

Provides an interface to call arbitrary class methods of vectorized environments

Parameters

- **method_name** – (str) The name of the env class method to invoke
- **method_args** – (tuple) Any positional arguments to provide in the call
- **method_kwargs** – (dict) Any keyword arguments to provide in the call

Returns (list) List of items returned by the environment's method call

get_attr (*attr_name*)

Provides a mechanism for getting class attributes from vectorized environments

Parameters `attr_name` – (str) The name of the attribute whose value to return

Returns (list) List of values of 'attr_name' in all environments

get_images ()

Return RGB images from each environment

render (**args*, ***kwargs*)

Gym environment rendering

Parameters `mode` – (str) the rendering type

reset ()

Reset all the environments and return an array of observations, or a tuple of observation arrays.

If `step_async` is still doing work, that work will be cancelled and `step_wait()` should not be called until `step_async()` is invoked again.

Returns ([int] or [float]) observation

set_attr (*attr_name*, *value*, *indices=None*)

Provides a mechanism for setting arbitrary class attributes inside vectorized environments

Parameters

- **attr_name** – (str) Name of attribute to assign new value
- **value** – (obj) Value to assign to 'attr_name'
- **indices** – (list,int) Indices of envs to assign value

Returns (list) in case env access methods might return something, they will be returned in a list

step_async (*actions*)

Tell all the environments to start taking a step with the given actions. Call `step_wait()` to get the results of the step.

You should not call this if a `step_async` run is already pending.

step_wait ()

Wait for the step taken with `step_async()`.

Returns ([int] or [float], [float], [bool], dict) observation, reward, done, information

1.5.2 SubprocVecEnv

class `stable_baselines.common.vec_env.SubprocVecEnv` (*env_fns*)

Creates a multiprocess vectorized wrapper for multiple environments

Parameters **env_fns** – ([Gym Environment]) Environments to run in subprocesses

close ()

Clean up the environment’s resources.

env_method (*method_name*, **method_args*, ***method_kwargs*)

Provides an interface to call arbitrary class methods of vectorized environments

Parameters

- **method_name** – (str) The name of the env class method to invoke
- **method_args** – (tuple) Any positional arguments to provide in the call
- **method_kwargs** – (dict) Any keyword arguments to provide in the call

Returns (list) List of items returned by each environment’s method call

get_attr (*attr_name*)

Provides a mechanism for getting class attributes from vectorized environments (note: attribute value returned must be picklable)

Parameters **attr_name** – (str) The name of the attribute whose value to return

Returns (list) List of values of ‘attr_name’ in all environments

get_images ()

Return RGB images from each environment

render (*mode*=‘human’, **args*, ***kwargs*)

Gym environment rendering

Parameters **mode** – (str) the rendering type

reset ()

Reset all the environments and return an array of observations, or a tuple of observation arrays.

If `step_async` is still doing work, that work will be cancelled and `step_wait()` should not be called until `step_async()` is invoked again.

Returns ([int] or [float]) observation

set_attr (*attr_name*, *value*, *indices*=None)

Provides a mechanism for setting arbitrary class attributes inside vectorized environments (note: this is a broadcast of a single value to all instances) (note: the value must be picklable)

Parameters

- **attr_name** – (str) Name of attribute to assign new value

- **value** – (obj) Value to assign to ‘attr_name’
- **indices** – (list,tuple) Iterable containing indices of envs whose attr to set

Returns (list) in case env access methods might return something, they will be returned in a list

step_async (*actions*)

Tell all the environments to start taking a step with the given actions. Call `step_wait()` to get the results of the step.

You should not call this if a `step_async` run is already pending.

step_wait ()

Wait for the step taken with `step_async()`.

Returns ([int] or [float], [float], [bool], dict) observation, reward, done, information

1.5.3 Wrappers

VecFrameStack

class `stable_baselines.common.vec_env.VecFrameStack` (*venv, n_stack*)

Frame stacking wrapper for vectorized environment

Parameters

- **venv** – (VecEnv) the vectorized environment to wrap
- **n_stack** – (int) Number of frames to stack

close ()

Clean up the environment’s resources.

reset ()

Reset all environments

step_wait ()

Wait for the step taken with `step_async()`.

Returns ([int] or [float], [float], [bool], dict) observation, reward, done, information

VecNormalize

class `stable_baselines.common.vec_env.VecNormalize` (*venv, training=True, norm_obs=True, norm_reward=True, clip_obs=10.0, clip_reward=10.0, gamma=0.99, epsilon=1e-08*)

A moving average, normalizing wrapper for vectorized environment. has support for saving/loading moving average,

Parameters

- **venv** – (VecEnv) the vectorized environment to wrap
- **training** – (bool) Whether to update or not the moving average
- **norm_obs** – (bool) Whether to normalize observation or not (default: True)
- **norm_reward** – (bool) Whether to normalize rewards or not (default: False)
- **clip_obs** – (float) Max absolute value for observation

- **clip_reward** – (float) Max value absolute for discounted reward
- **gamma** – (float) discount factor
- **epsilon** – (float) To avoid division by zero

get_original_obs()
returns the unnormalized observation

Returns (numpy float)

load_running_average(path)

Parameters path – (str) path to log dir

reset()
Reset all environments

save_running_average(path)

Parameters path – (str) path to log dir

step_wait()
Apply sequence of actions to sequence of environments actions -> (observations, rewards, news)
where ‘news’ is a boolean vector indicating whether each element is new.

1.6 Using Custom Environments

To use the rl baselines with custom environments, they just need to follow the *gym* interface. That is to say, your environment must implement the following methods (and inherits from OpenAI Gym Class):

```
import gym
from gym import spaces

class CustomEnv(gym.Env):
    """Custom Environment that follows gym interface"""
    metadata = {'render.modes': ['human']}

    def __init__(self, arg1, arg2, ...):
        super(CustomEnv, self).__init__()
        # Define action and observation space
        # They must be gym.spaces objects
        # Example when using discrete actions:
        self.action_space = spaces.Discrete(N_DISCRETE_ACTIONS)
        # Example for using image as input:
        self.observation_space = spaces.Box(low=0, high=255,
                                             shape=(HEIGHT, WIDTH, N_CHANNELS), dtype=np.
↳uint8)

    def step(self, action):
        ...
    def reset(self):
        ...
    def render(self, mode='human', close=False):
        ...
```

Then you can define and train a RL agent with:

```
# Instantiate and wrap the env
env = DummyVecEnv([lambda: CustomEnv(arg1, ...)])
# Define and Train the agent
model = A2C(CnnPolicy, env).learn(total_timesteps=1000)
```

You can find a [complete guide online](#) on creating a custom Gym environment.

Optionnaly, you can also register the environment with gym, that will allow you to create the RL agent in one line (and use `gym.make()` to instantiate the env).

In the project, for testing purposes, we use a custom environment named `IdentityEnv` defined [in this file](#). An example of how to use it can be found [here](#).

1.7 Custom Policy Network

Stable baselines provides default policy networks (see *Policies*) for images (CNNPolicies) and other type of input features (MlpPolicies). However, you can also easily define a custom architecture for the policy (or value) network:

```
import gym

from stable_baselines.common.policies import FeedForwardPolicy, register_policy
from stable_baselines.common.vec_env import DummyVecEnv
from stable_baselines import A2C

# Custom MLP policy of three layers of size 128 each
class CustomPolicy(FeedForwardPolicy):
    def __init__(self, *args, **kwargs):
        super(CustomPolicy, self).__init__(*args, **kwargs,
                                           layers=[128, 128, 128],
                                           feature_extraction="mlp")

# Create and wrap the environment
env = gym.make('LunarLander-v2')
env = DummyVecEnv([lambda: env])

model = A2C(CustomPolicy, env, verbose=1)
# Train the agent
model.learn(total_timesteps=100000)

del model
# When loading a model with a custom policy
# you MUST pass explicitly the policy when loading the saved model
model = A2C.load(policy=CustomPolicy)
```

Warning: When loading a model with a custom policy, you must pass the custom policy explicitly when loading the model. (cf previous example)

You can also registered your policy, to help with code simplicity: you can refer to your custom policy using a string.

```
import gym

from stable_baselines.common.policies import FeedForwardPolicy, register_policy
from stable_baselines.common.vec_env import DummyVecEnv
```

(continues on next page)

(continued from previous page)

```

from stable_baselines import A2C

# Custom MLP policy of three layers of size 128 each
class CustomPolicy(FeedForwardPolicy):
    def __init__(self, *args, **kwargs):
        super(CustomPolicy, self).__init__(*args, **kwargs,
                                           layers=[128, 128, 128],
                                           feature_extraction="mlp")

# Register the policy, it will check that the name is not already taken
register_policy('CustomPolicy', CustomPolicy)

# Because the policy is now registered, you can pass
# a string to the agent constructor instead of passing a class
model = A2C(policy='CustomPolicy', env='LunarLander-v2', verbose=1).learn(total_
↳ timesteps=100000)

```

If however, your task requires a more granular control over the policy architecture, you can redefine the policy directly:

```

import gym
import tensorflow as tf

from stable_baselines.common.policies import ActorCriticPolicy, register_policy, _
↳ nature_cnn
from stable_baselines.common.vec_env import DummyVecEnv
from stable_baselines import A2C

# Custom MLP policy of three layers of size 128 each for the actor and 2 layers of 32_
↳ for the critic,
# with a nature_cnn feature extractor
class CustomPolicy(ActorCriticPolicy):
    def __init__(self, sess, ob_space, ac_space, n_env, n_steps, n_batch, reuse=False,
↳ **kwargs):
        super(CustomPolicy, self).__init__(sess, ob_space, ac_space, n_env, n_steps, _
↳ n_batch, n_lstm=256,
                                           reuse=reuse, scale=True)

        with tf.variable_scope("model", reuse=reuse):
            activ = tf.nn.relu

            extracted_features = nature_cnn(self.processed_x, **kwargs)
            extracted_features = tf.layers.flatten(extracted_features)

            pi_h = extracted_features
            for i, layer_size in enumerate([128, 128, 128]):
                pi_h = activ(tf.layers.dense(pi_h, layer_size, name='pi_fc' + str(i)))
            pi_latent = pi_h

            vf_h = extracted_features
            for i, layer_size in enumerate([32, 32]):
                vf_h = activ(tf.layers.dense(vf_h, layer_size, name='vf_fc' + str(i)))
            value_fn = tf.layers.dense(vf_h, 1, name='vf')
            vf_latent = vf_h

            self.proba_distribution, self.policy, self.q_value = \
                self.pdtype.proba_distribution_from_latent(pi_latent, vf_latent, init_
↳ scale=0.01)

```

(continues on next page)

(continued from previous page)

```

        self.value_fn = value_fn
        self.initial_state = None
        self._setup_init()

    def step(self, obs, state=None, mask=None):
        action, value, neglogp = self.sess.run([self.action, self._value, self.
↪neglogp], {self.obs_ph: obs})
        return action, value, self.initial_state, neglogp

    def proba_step(self, obs, state=None, mask=None):
        return self.sess.run(self.policy_proba, {self.obs_ph: obs})

    def value(self, obs, state=None, mask=None):
        return self.sess.run(self._value, {self.obs_ph: obs})

# Create and wrap the environment
env = gym.make('Breakout-v0')
env = DummyVecEnv([lambda: env])

model = A2C(CustomPolicy, env, verbose=1)
# Train the agent
model.learn(total_timesteps=100000)

```

1.8 Tensorboard Integration

1.8.1 Basic Usage

To use Tensorboard with the rl baselines, you simply need to define a log location for the RL agent:

```

import gym

from stable_baselines.common.policies import MlpPolicy
from stable_baselines.common.vec_env import DummyVecEnv
from stable_baselines import A2C

env = gym.make('CartPole-v1')
env = DummyVecEnv([lambda: env]) # The algorithms require a vectorized environment,
↪to run

model = A2C(MlpPolicy, env, verbose=1, tensorboard_log="./a2c_cartpole_tensorboard/")
model.learn(total_timesteps=10000)

```

Or after loading an existing model (by default the log path is not saved):

```

import gym

from stable_baselines.common.vec_env import DummyVecEnv
from stable_baselines import A2C

env = gym.make('CartPole-v1')
env = DummyVecEnv([lambda: env]) # The algorithms require a vectorized environment,
↪to run

```

(continues on next page)

(continued from previous page)

```
model = A2C.load("./a2c_cartpole.pkl", env=env, tensorboard_log="./a2c_cartpole_
↳tensorboard/")
model.learn(total_timesteps=10000)
```

You can also define custom logging name when training (by default it is the algorithm name)

```
import gym

from stable_baselines.common.vec_env import DummyVecEnv
from stable_baselines.common.policies import MlpPolicy
from stable_baselines import A2C

env = gym.make('CartPole-v1')
env = DummyVecEnv([lambda: env]) # The algorithms require a vectorized environment,
↳to run

model = A2C(MlpPolicy, env, verbose=1, tensorboard_log="./a2c_cartpole_tensorboard/")
model.learn(total_timesteps=10000, tb_log_name="first_run")
model.learn(total_timesteps=10000, tb_log_name="second_run")
model.learn(total_timesteps=10000, tb_log_name="thrid_run")
```

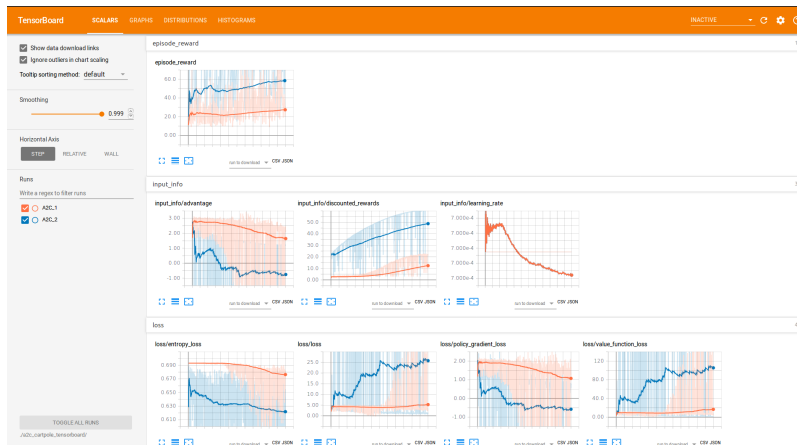
Once the learn function is called, you can monitor the RL agent during or after the training, with the following bash command:

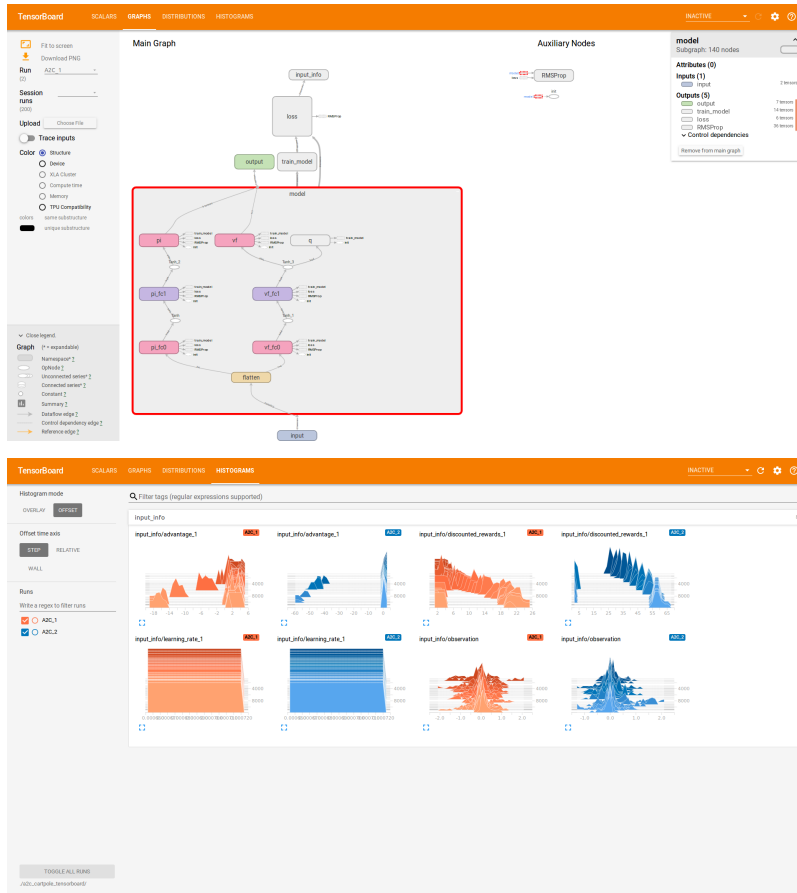
```
tensorboard --logdir ./a2c_cartpole_tensorboard/
```

you can also add past logging folders:

```
tensorboard --logdir ./a2c_cartpole_tensorboard/; ./ppo2_cartpole_tensorboard/
```

It will display information such as the model graph, the episode reward, the model losses, the observation and other parameter unique to some models.





1.8.2 Legacy Integration

All the information displayed in the terminal (default logging) can be also logged in tensorboard. For that, you need to define several environment variables:

```
# formats are comma-separated, but for tensorboard you only need the last one
# stdout -> terminal
export OPENAI_LOG_FORMAT='stdout,log,csv,tensorboard'
export OPENAI_LOGDIR=path/to/tensorboard/data
```

Then start tensorboard with:

```
tensorboard --logdir=$OPENAI_LOGDIR
```

1.9 Base RL Class

Common interface for all the RL algorithms

```
class stable_baselines.common.base_class.BaseRLModel (policy, env, verbose=0, *, re-
quires_vec_env, policy_base)
```

The base RL model

Parameters

- **policy** – (BasePolicy) Policy object
- **env** – (Gym environment) The environment to learn from (if registered in Gym, can be str. Can be None for loading trained models)
- **verbose** – (int) the verbosity level: 0 none, 1 training information, 2 tensorflow debug
- **requires_vec_env** – (bool) Does this model require a vectorized environment
- **policy_base** – (BasePolicy) the base policy used by this method

action_probability (*observation, state=None, mask=None*)
Get the model's action probability distribution from an observation

Parameters

- **observation** – (np.ndarray) the input observation
- **state** – (np.ndarray) The last states (can be None, used in recurrent policies)
- **mask** – (np.ndarray) The last masks (can be None, used in recurrent policies)

Returns (np.ndarray) the model's action probability distribution

get_env ()
returns the current environment (can be None if not defined)

Returns (Gym Environment) The current environment

learn (*total_timesteps, callback=None, seed=None, log_interval=100, tb_log_name='run'*)
Return a trained model.

Parameters

- **total_timesteps** – (int) The total number of samples to train on
- **seed** – (int) The initial seed for training, if None: keep current seed
- **callback** – (function (dict, dict)) function called at every steps with state of the algorithm. It takes the local and global variables.
- **log_interval** – (int) The number of timesteps before logging.
- **tb_log_name** – (str) the name of the run for tensorboard log

Returns (BaseRLModel) the trained model

classmethod load (*load_path, env=None, **kwargs*)
Load the model from file

Parameters

- **load_path** – (str) the saved parameter location
- **env** – (Gym Environment) the new environment to run the loaded model on (can be None if you only need prediction from a trained model)
- **kwargs** – extra arguments to change the model when loading

predict (*observation, state=None, mask=None, deterministic=False*)
Get the model's action from an observation

Parameters

- **observation** – (np.ndarray) the input observation
- **state** – (np.ndarray) The last states (can be None, used in recurrent policies)
- **mask** – (np.ndarray) The last masks (can be None, used in recurrent policies)

- **deterministic** – (bool) Whether or not to return deterministic actions.

Returns (np.ndarray, np.ndarray) the model’s action and the next state (used in recurrent policies)

save (*save_path*)

Save the current parameters to file

Parameters **save_path** – (str) the save location

set_env (*env*)

Checks the validity of the environment, and if it is coherent, set it as the current environment.

Parameters **env** – (Gym Environment) The environment for learning a policy

setup_model ()

Create all the functions and tensorflow graphs necessary to train the model

1.10 Policy Networks

Stable-baselines provides a set of default policies, that can be used with most action spaces. If you need more control on the policy architecture, You can also create a custom policy (see [Custom Policy Network](#)).

Note: CnnPolicies are for images only. MlpPolicies are made for other type of features (e.g. robot joints)

Warning: For all algorithms (except DDPG), continuous actions are only clipped during training (to avoid out of bound error). However, you have to manually clip the action when using the *predict()* method.

Available Policies

<i>MlpPolicy</i>	Policy object that implements actor critic, using a MLP (2 layers of 64)
<i>MlpLstmPolicy</i>	Policy object that implements actor critic, using LSTMs with a MLP feature extraction
<i>MlpLnLstmPolicy</i>	Policy object that implements actor critic, using a layer normalized LSTMs with a MLP feature extraction
<i>CnnPolicy</i>	Policy object that implements actor critic, using a CNN (the nature CNN)
<i>CnnLstmPolicy</i>	Policy object that implements actor critic, using LSTMs with a CNN feature extraction
<i>CnnLnLstmPolicy</i>	Policy object that implements actor critic, using a layer normalized LSTMs with a CNN feature extraction

1.10.1 Base Classes

```
class stable_baselines.common.policies.ActorCriticPolicy (sess, ob_space,
ac_space, n_env, n_steps,
n_batch, n_lstm=256,
reuse=False,
scale=False)
```

Policy object that implements actor critic

Parameters

- **sess** – (TensorFlow session) The current TensorFlow session
- **ob_space** – (Gym Space) The observation space of the environment
- **ac_space** – (Gym Space) The action space of the environment
- **n_env** – (int) The number of environments to run
- **n_steps** – (int) The number of steps to run for each environment
- **n_batch** – (int) The number of batch to run ($n_{\text{envs}} * n_{\text{steps}}$)
- **n_lstm** – (int) The number of LSTM cells (for recurrent policies)
- **reuse** – (bool) If the policy is reusable or not
- **scale** – (bool) whether or not to scale the input

```
proba_step (obs, state=None, mask=None)
```

Returns the action probability for a single step

Parameters

- **obs** – ([float] or [int]) The current observation of the environment
- **state** – ([float]) The last states (used in recurrent policies)
- **mask** – ([float]) The last masks (used in recurrent policies)

Returns ([float]) the action probability

```
step (obs, state=None, mask=None, deterministic=False)
```

Returns the policy for a single step

Parameters

- **obs** – ([float] or [int]) The current observation of the environment
- **state** – ([float]) The last states (used in recurrent policies)
- **mask** – ([float]) The last masks (used in recurrent policies)
- **deterministic** – (bool) Whether or not to return deterministic actions.

Returns ([float], [float], [float], [float]) actions, values, states, neglogp

```
value (obs, state=None, mask=None)
```

Returns the value for a single step

Parameters

- **obs** – ([float] or [int]) The current observation of the environment
- **state** – ([float]) The last states (used in recurrent policies)
- **mask** – ([float]) The last masks (used in recurrent policies)

Returns ([float]) The associated value of the action

```
class stable_baselines.common.policies.FeedForwardPolicy (sess, ob_space,
                                                         ac_space, n_env, n_steps,
                                                         n_batch, reuse=False,
                                                         layers=None,
                                                         cnn_extractor=<function
                                                         nature_cnn>, fea-
                                                         ture_extraction='cnn',
                                                         **kwargs)
```

Policy object that implements actor critic, using a feed forward neural network.

Parameters

- **sess** – (TensorFlow session) The current TensorFlow session
- **ob_space** – (Gym Space) The observation space of the environment
- **ac_space** – (Gym Space) The action space of the environment
- **n_env** – (int) The number of environments to run
- **n_steps** – (int) The number of steps to run for each environment
- **n_batch** – (int) The number of batch to run ($n_{\text{envs}} * n_{\text{steps}}$)
- **reuse** – (bool) If the policy is reusable or not
- **layers** – ([int]) The size of the Neural network for the policy (if None, default to [64, 64])
- **cnn_extractor** – (function (TensorFlow Tensor, ***kwargs*): (TensorFlow Tensor)) the CNN feature extraction
- **feature_extraction** – (str) The feature extraction type (“cnn” or “mlp”)
- **kwargs** – (dict) Extra keyword arguments for the nature CNN feature extraction

proba_step (*obs*, *state=None*, *mask=None*)

Returns the action probability for a single step

Parameters

- **obs** – ([float] or [int]) The current observation of the environment
- **state** – ([float]) The last states (used in recurrent policies)
- **mask** – ([float]) The last masks (used in recurrent policies)

Returns ([float]) the action probability

step (*obs*, *state=None*, *mask=None*, *deterministic=False*)

Returns the policy for a single step

Parameters

- **obs** – ([float] or [int]) The current observation of the environment
- **state** – ([float]) The last states (used in recurrent policies)
- **mask** – ([float]) The last masks (used in recurrent policies)
- **deterministic** – (bool) Whether or not to return deterministic actions.

Returns ([float], [float], [float], [float]) actions, values, states, neglogp

value (*obs*, *state=None*, *mask=None*)

Returns the value for a single step

Parameters

- **obs** – ([float] or [int]) The current observation of the environment
- **state** – ([float]) The last states (used in recurrent policies)
- **mask** – ([float]) The last masks (used in recurrent policies)

Returns ([float]) The associated value of the action

```
class stable_baselines.common.policies.LstmPolicy(sess, ob_space, ac_space, n_env,
                                                  n_steps, n_batch, n_lstm=256,
                                                  reuse=False, layers=None,
                                                  cnn_extractor=<function nature_cnn>, layer_norm=False, feature_extraction='cnn', **kwargs)
```

Policy object that implements actor critic, using LSTMs.

Parameters

- **sess** – (TensorFlow session) The current TensorFlow session
- **ob_space** – (Gym Space) The observation space of the environment
- **ac_space** – (Gym Space) The action space of the environment
- **n_env** – (int) The number of environments to run
- **n_steps** – (int) The number of steps to run for each environment
- **n_batch** – (int) The number of batch to run (n_envs * n_steps)
- **n_lstm** – (int) The number of LSTM cells (for recurrent policies)
- **reuse** – (bool) If the policy is reusable or not
- **layers** – ([int]) The size of the Neural network before the LSTM layer (if None, default to [64, 64])
- **cnn_extractor** – (function (TensorFlow Tensor, **kwargs): (TensorFlow Tensor)) the CNN feature extraction
- **layer_norm** – (bool) Whether or not to use layer normalizing LSTMs
- **feature_extraction** – (str) The feature extraction type (“cnn” or “mlp”)
- **kwargs** – (dict) Extra keyword arguments for the nature CNN feature extraction

proba_step (obs, state=None, mask=None)

Returns the action probability for a single step

Parameters

- **obs** – ([float] or [int]) The current observation of the environment
- **state** – ([float]) The last states (used in recurrent policies)
- **mask** – ([float]) The last masks (used in recurrent policies)

Returns ([float]) the action probability

step (obs, state=None, mask=None, deterministic=False)

Returns the policy for a single step

Parameters

- **obs** – ([float] or [int]) The current observation of the environment
- **state** – ([float]) The last states (used in recurrent policies)
- **mask** – ([float]) The last masks (used in recurrent policies)

- **deterministic** – (bool) Whether or not to return deterministic actions.

Returns ([float], [float], [float], [float]) actions, values, states, neglogp

value (*obs*, *state=None*, *mask=None*)

Returns the value for a single step

Parameters

- **obs** – ([float] or [int]) The current observation of the environment
- **state** – ([float]) The last states (used in recurrent policies)
- **mask** – ([float]) The last masks (used in recurrent policies)

Returns ([float]) The associated value of the action

1.10.2 MLP Policies

```
class stable_baselines.common.policies.MlpPolicy(sess, ob_space, ac_space, n_env,
                                                n_steps, n_batch, reuse=False,
                                                **kwargs)
```

Policy object that implements actor critic, using a MLP (2 layers of 64)

Parameters

- **sess** – (TensorFlow session) The current TensorFlow session
- **ob_space** – (Gym Space) The observation space of the environment
- **ac_space** – (Gym Space) The action space of the environment
- **n_env** – (int) The number of environments to run
- **n_steps** – (int) The number of steps to run for each environment
- **n_batch** – (int) The number of batch to run ($n_{\text{envs}} * n_{\text{steps}}$)
- **reuse** – (bool) If the policy is reusable or not
- **_kwargs** – (dict) Extra keyword arguments for the nature CNN feature extraction

```
class stable_baselines.common.policies.MlpLstmPolicy(sess, ob_space, ac_space,
                                                    n_env, n_steps, n_batch,
                                                    n_lstm=256, reuse=False,
                                                    **kwargs)
```

Policy object that implements actor critic, using LSTMs with a MLP feature extraction

Parameters

- **sess** – (TensorFlow session) The current TensorFlow session
- **ob_space** – (Gym Space) The observation space of the environment
- **ac_space** – (Gym Space) The action space of the environment
- **n_env** – (int) The number of environments to run
- **n_steps** – (int) The number of steps to run for each environment
- **n_batch** – (int) The number of batch to run ($n_{\text{envs}} * n_{\text{steps}}$)
- **n_lstm** – (int) The number of LSTM cells (for recurrent policies)
- **reuse** – (bool) If the policy is reusable or not
- **kwargs** – (dict) Extra keyword arguments for the nature CNN feature extraction

```
class stable_baselines.common.policies.MlpLnLstmPolicy(sess, ob_space, ac_space,
                                                    n_env, n_steps, n_batch,
                                                    n_lstm=256, reuse=False,
                                                    **kwargs)
```

Policy object that implements actor critic, using a layer normalized LSTMs with a MLP feature extraction

Parameters

- **sess** – (TensorFlow session) The current TensorFlow session
- **ob_space** – (Gym Space) The observation space of the environment
- **ac_space** – (Gym Space) The action space of the environment
- **n_env** – (int) The number of environments to run
- **n_steps** – (int) The number of steps to run for each environment
- **n_batch** – (int) The number of batch to run ($n_{\text{envs}} * n_{\text{steps}}$)
- **n_lstm** – (int) The number of LSTM cells (for recurrent policies)
- **reuse** – (bool) If the policy is reusable or not
- **kwargs** – (dict) Extra keyword arguments for the nature CNN feature extraction

1.10.3 CNN Policies

```
class stable_baselines.common.policies.CnnPolicy(sess, ob_space, ac_space, n_env,
                                                    n_steps, n_batch, reuse=False,
                                                    **kwargs)
```

Policy object that implements actor critic, using a CNN (the nature CNN)

Parameters

- **sess** – (TensorFlow session) The current TensorFlow session
- **ob_space** – (Gym Space) The observation space of the environment
- **ac_space** – (Gym Space) The action space of the environment
- **n_env** – (int) The number of environments to run
- **n_steps** – (int) The number of steps to run for each environment
- **n_batch** – (int) The number of batch to run ($n_{\text{envs}} * n_{\text{steps}}$)
- **reuse** – (bool) If the policy is reusable or not
- **kwargs** – (dict) Extra keyword arguments for the nature CNN feature extraction

```
class stable_baselines.common.policies.CnnLstmPolicy(sess, ob_space, ac_space,
                                                    n_env, n_steps, n_batch,
                                                    n_lstm=256, reuse=False,
                                                    **kwargs)
```

Policy object that implements actor critic, using LSTMs with a CNN feature extraction

Parameters

- **sess** – (TensorFlow session) The current TensorFlow session
- **ob_space** – (Gym Space) The observation space of the environment
- **ac_space** – (Gym Space) The action space of the environment
- **n_env** – (int) The number of environments to run

- **n_steps** – (int) The number of steps to run for each environment
- **n_batch** – (int) The number of batch to run ($n_{\text{envs}} * n_{\text{steps}}$)
- **n_lstm** – (int) The number of LSTM cells (for recurrent policies)
- **reuse** – (bool) If the policy is reusable or not
- **kwargs** – (dict) Extra keyword arguments for the nature CNN feature extraction

```
class stable_baselines.common.policies.CnnLnLstmPolicy(sess, ob_space, ac_space,
                                                    n_env, n_steps, n_batch,
                                                    n_lstm=256, reuse=False,
                                                    **kwargs)
```

Policy object that implements actor critic, using a layer normalized LSTMs with a CNN feature extraction

Parameters

- **sess** – (TensorFlow session) The current TensorFlow session
- **ob_space** – (Gym Space) The observation space of the environment
- **ac_space** – (Gym Space) The action space of the environment
- **n_env** – (int) The number of environments to run
- **n_steps** – (int) The number of steps to run for each environment
- **n_batch** – (int) The number of batch to run ($n_{\text{envs}} * n_{\text{steps}}$)
- **n_lstm** – (int) The number of LSTM cells (for recurrent policies)
- **reuse** – (bool) If the policy is reusable or not
- **kwargs** – (dict) Extra keyword arguments for the nature CNN feature extraction

1.11 A2C

A synchronous, deterministic variant of [Asynchronous Advantage Actor Critic \(A3C\)](#). It uses multiple workers to avoid the use of a replay buffer.

1.11.1 Notes

- Original paper: <https://arxiv.org/abs/1602.01783>
- OpenAI blog post: <https://blog.openai.com/openai-baselines-ppo/>
- `python -m stable_baselines.ppo2.run_atari` runs the algorithm for 40M frames = 10M timesteps on an Atari game. See help (-h) for more options.
- `python -m stable_baselines.ppo2.run_mujoco` runs the algorithm for 1M frames on a Mujoco environment.

1.11.2 Can I use?

- Recurrent policies:
- Multi processing:
- Gym spaces:

Space	Action	Observation
Discrete		
Box		
MultiDiscrete		
MultiBinary		

1.11.3 Example

Train a A2C agent on *CartPole-v1* using 4 processes.

```
import gym

from stable_baselines.common.policies import MlpPolicy
from stable_baselines.common.vec_env import SubprocVecEnv
from stable_baselines import A2C

# multiprocessing environment
n_cpu = 4
env = SubprocVecEnv([lambda: gym.make('CartPole-v1') for i in range(n_cpu)])

model = A2C(MlpPolicy, env, verbose=1)
model.learn(total_timesteps=25000)
model.save("a2c_cartpole")

del model # remove to demonstrate saving and loading

model = A2C.load("a2c_cartpole")

obs = env.reset()
while True:
    action, _states = model.predict(obs)
    obs, rewards, dones, info = env.step(action)
    env.render()
```

1.11.4 Parameters

```
class stable_baselines.a2c.A2C(policy, env, gamma=0.99, n_steps=5, vf_coef=0.25,
                               ent_coef=0.01, max_grad_norm=0.5, learning_rate=0.0007,
                               alpha=0.99, epsilon=1e-05, lr_schedule='linear', verbose=0,
                               tensorboard_log=None, _init_setup_model=True)
```

The A2C (Advantage Actor Critic) model class, <https://arxiv.org/abs/1602.01783>

Parameters

- **policy** – (ActorCriticPolicy or str) The policy model to use (MlpPolicy, CnnPolicy, CnnLstmPolicy, ...)
- **env** – (Gym environment or str) The environment to learn from (if registered in Gym, can be str)
- **gamma** – (float) Discount factor
- **n_steps** – (int) The number of steps to run for each environment per update (i.e. batch size is $n_steps * n_env$ where n_env is number of environment copies running in parallel)
- **vf_coef** – (float) Value function coefficient for the loss calculation

- **ent_coef** – (float) Entropy coefficient for the loss calculation
- **max_grad_norm** – (float) The maximum value for the gradient clipping
- **learning_rate** – (float) The learning rate
- **alpha** – (float) RMSProp decay parameter (default: 0.99)
- **epsilon** – (float) RMSProp epsilon (stabilizes square root computation in denominator of RMSProp update) (default: 1e-5)
- **lr_schedule** – (str) The type of scheduler for the learning rate update ('linear', 'constant', 'double_linear_con', 'middle_drop' or 'double_middle_drop')
- **verbose** – (int) the verbosity level: 0 none, 1 training information, 2 tensorflow debug
- **tensorboard_log** – (str) the log location for tensorboard (if None, no logging)
- **_init_setup_model** – (bool) Whether or not to build the network at the creation of the instance (used only for loading)

action_probability (*observation, state=None, mask=None*)
Get the model's action probability distribution from an observation

Parameters

- **observation** – (np.ndarray) the input observation
- **state** – (np.ndarray) The last states (can be None, used in recurrent policies)
- **mask** – (np.ndarray) The last masks (can be None, used in recurrent policies)

Returns (np.ndarray) the model's action probability distribution

get_env ()
returns the current environment (can be None if not defined)

Returns (Gym Environment) The current environment

learn (*total_timesteps, callback=None, seed=None, log_interval=100, tb_log_name='A2C'*)
Return a trained model.

Parameters

- **total_timesteps** – (int) The total number of samples to train on
- **seed** – (int) The initial seed for training, if None: keep current seed
- **callback** – (function (dict, dict)) function called at every steps with state of the algorithm. It takes the local and global variables.
- **log_interval** – (int) The number of timesteps before logging.
- **tb_log_name** – (str) the name of the run for tensorboard log

Returns (BaseRLModel) the trained model

classmethod load (*load_path, env=None, **kwargs*)
Load the model from file

Parameters

- **load_path** – (str) the saved parameter location
- **env** – (Gym Environment) the new environment to run the loaded model on (can be None if you only need prediction from a trained model)
- **kwargs** – extra arguments to change the model when loading

predict (*observation*, *state=None*, *mask=None*, *deterministic=False*)

Get the model's action from an observation

Parameters

- **observation** – (np.ndarray) the input observation
- **state** – (np.ndarray) The last states (can be None, used in recurrent policies)
- **mask** – (np.ndarray) The last masks (can be None, used in recurrent policies)
- **deterministic** – (bool) Whether or not to return deterministic actions.

Returns (np.ndarray, np.ndarray) the model's action and the next state (used in recurrent policies)

save (*save_path*)

Save the current parameters to file

Parameters **save_path** – (str) the save location

set_env (*env*)

Checks the validity of the environment, and if it is coherent, set it as the current environment.

Parameters **env** – (Gym Environment) The environment for learning a policy

setup_model ()

Create all the functions and tensorflow graphs necessary to train the model

1.12 ACER

Sample Efficient Actor-Critic with Experience Replay (ACER) combines several ideas of previous algorithms: it uses multiple workers (as A2C), implements a replay buffer (as in DQN), uses Retrace for Q-value estimation, importance sampling and a trust region.

1.12.1 Notes

- Original paper: <https://arxiv.org/abs/1611.01224>
- `python -m stable_baselines.acer.run_atari` runs the algorithm for 40M frames = 10M timesteps on an Atari game. See help (-h) for more options.

1.12.2 Can I use?

- Recurrent policies:
- Multi processing:
- Gym spaces:

Space	Action	Observation
Discrete		
Box		
MultiDiscrete		
MultiBinary		

1.12.3 Example

```
import gym

from stable_baselines.common.policies import MlpPolicy, MlpLstmPolicy, MlpLnLstmPolicy
from stable_baselines.common.vec_env import SubprocVecEnv
from stable_baselines import ACER

# multiprocessing environment
n_cpu = 4
env = SubprocVecEnv([lambda: gym.make('CartPole-v1') for i in range(n_cpu)])

model = ACER(MlpPolicy, env, verbose=1)
model.learn(total_timesteps=25000)
model.save("acer_cartpole")

del model # remove to demonstrate saving and loading

model = ACER.load("acer_cartpole")

obs = env.reset()
while True:
    action, _states = model.predict(obs)
    obs, rewards, dones, info = env.step(action)
    env.render()
```

1.12.4 Parameters

```
class stable_baselines.acer.ACER(policy, env, gamma=0.99, n_steps=20, num_procs=1,
                                q_coef=0.5, ent_coef=0.01, max_grad_norm=10, learning_rate=0.0007, lr_schedule='linear', rprop_alpha=0.99,
                                rprop_epsilon=1e-05, buffer_size=5000, replay_ratio=4, replay_start=1000, correction_term=10.0, trust_region=True,
                                alpha=0.99, delta=1, verbose=0, tensorboard_log=None, _init_setup_model=True)
```

The ACER (Actor-Critic with Experience Replay) model class, <https://arxiv.org/abs/1611.01224>

Parameters

- **policy** – (ActorCriticPolicy or str) The policy model to use (MlpPolicy, CnnPolicy, CnnLstmPolicy, ...)
- **env** – (Gym environment or str) The environment to learn from (if registered in Gym, can be str)
- **gamma** – (float) The discount value
- **n_steps** – (int) The number of steps to run for each environment per update (i.e. batch size is $n_steps * n_env$ where n_env is number of environment copies running in parallel)
- **num_procs** – (int) The number of threads for TensorFlow operations
- **q_coef** – (float) The weight for the loss on the Q value
- **ent_coef** – (float) The weight for the entropic loss
- **max_grad_norm** – (float) The clipping value for the maximum gradient
- **learning_rate** – (float) The initial learning rate for the RMS prop optimizer

- **lr_schedule** – (str) The type of scheduler for the learning rate update ('linear', 'constant', 'double_linear_con', 'middle_drop' or 'double_middle_drop')
- **rprop_epsilon** – (float) RMSProp epsilon (stabilizes square root computation in denominator of RMSProp update) (default: 1e-5)
- **rprop_alpha** – (float) RMSProp decay parameter (default: 0.99)
- **buffer_size** – (int) The buffer size in number of steps
- **replay_ratio** – (float) The number of replay learning per on policy learning on average, using a poisson distribution
- **replay_start** – (int) The minimum number of steps in the buffer, before learning replay
- **correction_term** – (float) Importance weight clipping factor (default: 10)
- **trust_region** – (bool) Whether or not algorithms estimates the gradient KL divergence between the old and updated policy and uses it to determine step size (default: True)
- **alpha** – (float) The decay rate for the Exponential moving average of the parameters
- **delta** – (float) max KL divergence between the old policy and updated policy (default: 1)
- **verbose** – (int) the verbosity level: 0 none, 1 training information, 2 tensorflow debug
- **tensorboard_log** – (str) the log location for tensorboard (if None, no logging)
- **_init_setup_model** – (bool) Whether or not to build the network at the creation of the instance

action_probability (*observation, state=None, mask=None*)

Get the model's action probability distribution from an observation

Parameters

- **observation** – (np.ndarray) the input observation
- **state** – (np.ndarray) The last states (can be None, used in recurrent policies)
- **mask** – (np.ndarray) The last masks (can be None, used in recurrent policies)

Returns (np.ndarray) the model's action probability distribution

get_env ()

returns the current environment (can be None if not defined)

Returns (Gym Environment) The current environment

learn (*total_timesteps, callback=None, seed=None, log_interval=100, tb_log_name='ACER'*)

Return a trained model.

Parameters

- **total_timesteps** – (int) The total number of samples to train on
- **seed** – (int) The initial seed for training, if None: keep current seed
- **callback** – (function (dict, dict)) function called at every steps with state of the algorithm. It takes the local and global variables.
- **log_interval** – (int) The number of timesteps before logging.
- **tb_log_name** – (str) the name of the run for tensorboard log

Returns (BaseRLModel) the trained model

classmethod `load (load_path, env=None, **kwargs)`

Load the model from file

Parameters

- **load_path** – (str) the saved parameter location
- **env** – (Gym Environment) the new environment to run the loaded model on (can be None if you only need prediction from a trained model)
- **kwargs** – extra arguments to change the model when loading

predict (*observation, state=None, mask=None, deterministic=False*)

Get the model’s action from an observation

Parameters

- **observation** – (np.ndarray) the input observation
- **state** – (np.ndarray) The last states (can be None, used in recurrent policies)
- **mask** – (np.ndarray) The last masks (can be None, used in recurrent policies)
- **deterministic** – (bool) Whether or not to return deterministic actions.

Returns (np.ndarray, np.ndarray) the model’s action and the next state (used in recurrent policies)

save (*save_path*)

Save the current parameters to file

Parameters **save_path** – (str) the save location

set_env (*env*)

Checks the validity of the environment, and if it is coherent, set it as the current environment.

Parameters **env** – (Gym Environment) The environment for learning a policy

setup_model ()

Create all the functions and tensorflow graphs necessary to train the model

1.13 ACKTR

Actor Critic using Kronecker-Factored Trust Region (ACKTR) uses Kronecker-factored approximate curvature (KFAC) for trust region optimization.

1.13.1 Notes

- Original paper: <https://arxiv.org/abs/1708.05144>
- Baselines blog post: <https://blog.openai.com/baselines-acktr-a2c/>
- `python -m stable_baselines.acktr.run_atari` runs the algorithm for 40M frames = 10M timesteps on an Atari game. See help (-h) for more options.

1.13.2 Can I use?

- Recurrent policies:
- Multi processing:

- Gym spaces:

Space	Action	Observation
Discrete		
Box		
MultiDiscrete		
MultiBinary		

1.13.3 Example

```
import gym

from stable_baselines.common.policies import MlpPolicy, MlpLstmPolicy, MlpLnLstmPolicy
from stable_baselines.common.vec_env import SubprocVecEnv
from stable_baselines import ACKTR

# multiprocessing environment
n_cpu = 4
env = SubprocVecEnv([lambda: gym.make('CartPole-v1') for i in range(n_cpu)])

model = ACKTR(MlpPolicy, env, verbose=1)
model.learn(total_timesteps=25000)
model.save("acktr_cartpole")

del model # remove to demonstrate saving and loading

model = ACKTR.load("acktr_cartpole")

obs = env.reset()
while True:
    action, _states = model.predict(obs)
    obs, rewards, dones, info = env.step(action)
    env.render()
```

1.13.4 Parameters

```
class stable_baselines.acktr.ACKTR(policy, env, gamma=0.99, nprocs=1, n_steps=20,
                                   ent_coef=0.01, vf_coef=0.25, vf_fisher_coef=1.0, learning_rate=0.25,
                                   max_grad_norm=0.5, kfac_clip=0.001, lr_schedule='linear', verbose=0,
                                   tensorboard_log=None, _init_setup_model=True, async_eigen_decomp=False)
```

The ACKTR (Actor Critic using Kronecker-Factored Trust Region) model class, <https://arxiv.org/abs/1708.05144>

Parameters

- **policy** – (ActorCriticPolicy or str) The policy model to use (MlpPolicy, CnnPolicy, CnnLstmPolicy, ...)
- **env** – (Gym environment or str) The environment to learn from (if registered in Gym, can be str)
- **gamma** – (float) Discount factor
- **nprocs** – (int) The number of threads for TensorFlow operations

- **n_steps** – (int) The number of steps to run for each environment
- **ent_coef** – (float) The weight for the entropic loss
- **vf_coef** – (float) The weight for the loss on the value function
- **vf_fisher_coef** – (float) The weight for the fisher loss on the value function
- **learning_rate** – (float) The initial learning rate for the RMS prop optimizer
- **max_grad_norm** – (float) The clipping value for the maximum gradient
- **kfac_clip** – (float) gradient clipping for Kullback leiber
- **lr_schedule** – (str) The type of scheduler for the learning rate update ('linear', 'constant', 'double_linear_con', 'middle_drop' or 'double_middle_drop')
- **verbose** – (int) the verbosity level: 0 none, 1 training information, 2 tensorflow debug
- **tensorboard_log** – (str) the log location for tensorboard (if None, no logging)
- **_init_setup_model** – (bool) Whether or not to build the network at the creation of the instance
- **async_eigen_decomp** – (bool) Use async eigen decomposition

action_probability (*observation, state=None, mask=None*)

Get the model's action probability distribution from an observation

Parameters

- **observation** – (np.ndarray) the input observation
- **state** – (np.ndarray) The last states (can be None, used in recurrent policies)
- **mask** – (np.ndarray) The last masks (can be None, used in recurrent policies)

Returns (np.ndarray) the model's action probability distribution

get_env ()

returns the current environment (can be None if not defined)

Returns (Gym Environment) The current environment

learn (*total_timesteps, callback=None, seed=None, log_interval=100, tb_log_name='ACKTR'*)

Return a trained model.

Parameters

- **total_timesteps** – (int) The total number of samples to train on
- **seed** – (int) The initial seed for training, if None: keep current seed
- **callback** – (function (dict, dict)) function called at every steps with state of the algorithm. It takes the local and global variables.
- **log_interval** – (int) The number of timesteps before logging.
- **tb_log_name** – (str) the name of the run for tensorboard log

Returns (BaseRLModel) the trained model

classmethod load (*load_path, env=None, **kwargs*)

Load the model from file

Parameters

- **load_path** – (str) the saved parameter location

- **env** – (Gym Environment) the new environment to run the loaded model on (can be None if you only need prediction from a trained model)
- **kwargs** – extra arguments to change the model when loading

predict (*observation, state=None, mask=None, deterministic=False*)

Get the model’s action from an observation

Parameters

- **observation** – (np.ndarray) the input observation
- **state** – (np.ndarray) The last states (can be None, used in recurrent policies)
- **mask** – (np.ndarray) The last masks (can be None, used in recurrent policies)
- **deterministic** – (bool) Whether or not to return deterministic actions.

Returns (np.ndarray, np.ndarray) the model’s action and the next state (used in recurrent policies)

save (*save_path*)

Save the current parameters to file

Parameters **save_path** – (str) the save location

set_env (*env*)

Checks the validity of the environment, and if it is coherent, set it as the current environment.

Parameters **env** – (Gym Environment) The environment for learning a policy

setup_model ()

Create all the functions and tensorflow graphs necessary to train the model

1.14 DDPG

Deep Deterministic Policy Gradient (DDPG)

Warning: The DDPG model does not support `stable_baselines.common.policies` because it uses q-value instead of value estimation, as a result it must use its own policy models (see *DDPG Policies*).

Available Policies

<i>MlpPolicy</i>	Policy object that implements actor critic, using a MLP (2 layers of 64)
<i>LnMlpPolicy</i>	Policy object that implements actor critic, using a MLP (2 layers of 64), with layer normalisation
<i>CnnPolicy</i>	Policy object that implements actor critic, using a CNN (the nature CNN)
<i>LnCnnPolicy</i>	Policy object that implements actor critic, using a CNN (the nature CNN), with layer normalisation

1.14.1 Notes

- Original paper: <https://arxiv.org/abs/1509.02971>

- Baselines post: <https://blog.openai.com/better-exploration-with-parameter-noise/>
- `python -m stable_baselines.ddpg.main` runs the algorithm for 1M frames = 10M timesteps on a Mujoco environment. See help (-h) for more options.

1.14.2 Can I use?

- Recurrent policies:
- Multi processing:
- Gym spaces:

Space	Action	Observation
Discrete		
Box		
MultiDiscrete		
MultiBinary		

1.14.3 Example

```
import gym
import numpy as np

from stable_baselines.ddpg.policies import MlpPolicy
from stable_baselines.common.vec_env import DummyVecEnv
from stable_baselines.ddpg.noise import NormalActionNoise, OrnsteinUhlenbeckActionNoise, AdaptiveParamNoiseSpec
from stable_baselines import DDPG

env = gym.make('MountainCarContinuous-v0')
env = DummyVecEnv([lambda: env])

# the noise objects for DDPG
n_actions = env.action_space.shape[-1]
param_noise = None
action_noise = OrnsteinUhlenbeckActionNoise(mean=np.zeros(n_actions), sigma=float(0.5) * np.ones(n_actions))

model = DDPG(MlpPolicy, env, verbose=1, param_noise=param_noise, action_noise=action_noise)
model.learn(total_timesteps=400000)
model.save("ddpg_mountain")

del model # remove to demonstrate saving and loading

model = DDPG.load("ddpg_mountain")

obs = env.reset()
while True:
    action, _states = model.predict(obs)
    obs, rewards, dones, info = env.step(action)
    env.render()
```

1.14.4 Parameters

```
class stable_baselines.ddpg.DDPG(policy, env, gamma=0.99, memory_policy=None,
                                eval_env=None, nb_train_steps=50,
                                nb_rollout_steps=100, nb_eval_steps=100,
                                param_noise=None, action_noise=None, normal-
                                ize_observations=False, tau=0.001, batch_size=128,
                                param_noise_adaption_interval=50, nor-
                                malize_returns=False, enable_popart=False,
                                observation_range=(-5.0, 5.0), critic_l2_reg=0.0,
                                return_range=(-inf, inf), actor_lr=0.0001, critic_lr=0.001,
                                clip_norm=None, reward_scale=1.0, render=False, ren-
                                der_eval=False, memory_limit=100, verbose=0, tensor-
                                board_log=None, _init_setup_model=True)
```

Deep Deterministic Policy Gradient (DDPG) model

DDPG: <https://arxiv.org/pdf/1509.02971.pdf>

Parameters

- **policy** – (DDPGPolicy or str) The policy model to use (MlpPolicy, CnnPolicy, LnMlpPolicy, ...)
- **env** – (Gym environment or str) The environment to learn from (if registered in Gym, can be str)
- **gamma** – (float) the discount rate
- **memory_policy** – (Memory) the replay buffer (if None, default to `baselines.ddpg.memory.Memory`)
- **eval_env** – (Gym Environment) the evaluation environment (can be None)
- **nb_train_steps** – (int) the number of training steps
- **nb_rollout_steps** – (int) the number of rollout steps
- **nb_eval_steps** – (int) the number of evaluation steps
- **param_noise** – (AdaptiveParamNoiseSpec) the parameter noise type (can be None)
- **action_noise** – (ActionNoise) the action noise type (can be None)
- **param_noise_adaption_interval** – (int) apply param noise every N steps
- **tau** – (float) the soft update coefficient (keep old values, between 0 and 1)
- **normalize_returns** – (bool) should the critic output be normalized
- **enable_popart** – (bool) enable pop-art normalization of the critic output (<https://arxiv.org/pdf/1602.07714.pdf>)
- **normalize_observations** – (bool) should the observation be normalized
- **batch_size** – (int) the size of the batch for learning the policy
- **observation_range** – (tuple) the bounding values for the observation
- **return_range** – (tuple) the bounding values for the critic output
- **critic_l2_reg** – (float) l2 regularizer coefficient
- **actor_lr** – (float) the actor learning rate
- **critic_lr** – (float) the critic learning rate

- **clip_norm** – (float) clip the gradients (disabled if None)
- **reward_scale** – (float) the value the reward should be scaled by
- **render** – (bool) enable rendering of the environment
- **render_eval** – (bool) enable rendering of the evaluation environment
- **memory_limit** – (int) the max number of transitions to store
- **verbose** – (int) the verbosity level: 0 none, 1 training information, 2 tensorflow debug
- **tensorboard_log** – (str) the log location for tensorboard (if None, no logging)
- **_init_setup_model** – (bool) Whether or not to build the network at the creation of the instance

action_probability (*observation, state=None, mask=None*)

Get the model's action probability distribution from an observation

Parameters

- **observation** – (np.ndarray) the input observation
- **state** – (np.ndarray) The last states (can be None, used in recurrent policies)
- **mask** – (np.ndarray) The last masks (can be None, used in recurrent policies)

Returns (np.ndarray) the model's action probability distribution

get_env ()

returns the current environment (can be None if not defined)

Returns (Gym Environment) The current environment

learn (*total_timesteps, callback=None, seed=None, log_interval=100, tb_log_name='DDPG'*)

Return a trained model.

Parameters

- **total_timesteps** – (int) The total number of samples to train on
- **seed** – (int) The initial seed for training, if None: keep current seed
- **callback** – (function (dict, dict)) function called at every steps with state of the algorithm. It takes the local and global variables.
- **log_interval** – (int) The number of timesteps before logging.
- **tb_log_name** – (str) the name of the run for tensorboard log

Returns (BaseRLModel) the trained model

classmethod load (*load_path, env=None, **kwargs*)

Load the model from file

Parameters

- **load_path** – (str) the saved parameter location
- **env** – (Gym Environment) the new environment to run the loaded model on (can be None if you only need prediction from a trained model)
- **kwargs** – extra arguments to change the model when loading

predict (*observation, state=None, mask=None, deterministic=True*)

Get the model's action from an observation

Parameters

- **observation** – (np.ndarray) the input observation
- **state** – (np.ndarray) The last states (can be None, used in recurrent policies)
- **mask** – (np.ndarray) The last masks (can be None, used in recurrent policies)
- **deterministic** – (bool) Whether or not to return deterministic actions.

Returns (np.ndarray, np.ndarray) the model’s action and the next state (used in recurrent policies)

save (*save_path*)

Save the current parameters to file

Parameters **save_path** – (str) the save location

set_env (*env*)

Checks the validity of the environment, and if it is coherent, set it as the current environment.

Parameters **env** – (Gym Environment) The environment for learning a policy

setup_model ()

Create all the functions and tensorflow graphs necessary to train the model

1.14.5 DDPG Policies

class `stable_baselines.ddpg.MlpPolicy` (*sess, ob_space, ac_space, n_env, n_steps, n_batch, reuse=False, **kwargs*)

Policy object that implements actor critic, using a MLP (2 layers of 64)

Parameters

- **sess** – (TensorFlow session) The current TensorFlow session
- **ob_space** – (Gym Space) The observation space of the environment
- **ac_space** – (Gym Space) The action space of the environment
- **n_env** – (int) The number of environments to run
- **n_steps** – (int) The number of steps to run for each environment
- **n_batch** – (int) The number of batch to run ($n_{\text{envs}} * n_{\text{steps}}$)
- **reuse** – (bool) If the policy is reusable or not
- **_kwargs** – (dict) Extra keyword arguments for the nature CNN feature extraction

make_actor (*obs=None, reuse=False, scope='pi'*)

creates an actor object

Parameters

- **obs** – (TensorFlow Tensor) The observation placeholder (can be None for default placeholder)
- **reuse** – (bool) whether or not to reuse parameters
- **scope** – (str) the scope name of the actor

Returns (TensorFlow Tensor) the output tensor

make_critic (*obs=None, action=None, reuse=False, scope='qf'*)

creates a critic object

Parameters

- **obs** – (TensorFlow Tensor) The observation placeholder (can be None for default placeholder)
- **action** – (TensorFlow Tensor) The action placeholder (can be None for default placeholder)
- **reuse** – (bool) whether or not to reuse parameters
- **scope** – (str) the scope name of the critic

Returns (TensorFlow Tensor) the output tensor

proba_step (*obs, state=None, mask=None*)

Returns the action probability for a single step

Parameters

- **obs** – ([float] or [int]) The current observation of the environment
- **state** – ([float]) The last states (used in recurrent policies)
- **mask** – ([float]) The last masks (used in recurrent policies)

Returns ([float]) the action probability

step (*obs, state=None, mask=None*)

Returns the policy for a single step

Parameters

- **obs** – ([float] or [int]) The current observation of the environment
- **state** – ([float]) The last states (used in recurrent policies)
- **mask** – ([float]) The last masks (used in recurrent policies)

Returns ([float]) actions

value (*obs, action, state=None, mask=None*)

Returns the value for a single step

Parameters

- **obs** – ([float] or [int]) The current observation of the environment
- **action** – ([float] or [int]) The taken action
- **state** – ([float]) The last states (used in recurrent policies)
- **mask** – ([float]) The last masks (used in recurrent policies)

Returns ([float]) The associated value of the action

class stable_baselines.ddpg.**LnMlpPolicy** (*sess, ob_space, ac_space, n_env, n_steps, n_batch, reuse=False, **kwargs*)

Policy object that implements actor critic, using a MLP (2 layers of 64), with layer normalisation

Parameters

- **sess** – (TensorFlow session) The current TensorFlow session
- **ob_space** – (Gym Space) The observation space of the environment
- **ac_space** – (Gym Space) The action space of the environment
- **n_env** – (int) The number of environments to run
- **n_steps** – (int) The number of steps to run for each environment
- **n_batch** – (int) The number of batch to run ($n_{\text{envs}} * n_{\text{steps}}$)

- **reuse** – (bool) If the policy is reusable or not
- **_kwargs** – (dict) Extra keyword arguments for the nature CNN feature extraction

make_actor (*obs=None, reuse=False, scope='pi'*)
creates an actor object

Parameters

- **obs** – (TensorFlow Tensor) The observation placeholder (can be None for default placeholder)
- **reuse** – (bool) whether or not to reuse parameters
- **scope** – (str) the scope name of the actor

Returns (TensorFlow Tensor) the output tensor

make_critic (*obs=None, action=None, reuse=False, scope='qf'*)
creates a critic object

Parameters

- **obs** – (TensorFlow Tensor) The observation placeholder (can be None for default placeholder)
- **action** – (TensorFlow Tensor) The action placeholder (can be None for default placeholder)
- **reuse** – (bool) whether or not to reuse parameters
- **scope** – (str) the scope name of the critic

Returns (TensorFlow Tensor) the output tensor

proba_step (*obs, state=None, mask=None*)
Returns the action probability for a single step

Parameters

- **obs** – ([float] or [int]) The current observation of the environment
- **state** – ([float]) The last states (used in recurrent policies)
- **mask** – ([float]) The last masks (used in recurrent policies)

Returns ([float]) the action probability

step (*obs, state=None, mask=None*)
Returns the policy for a single step

Parameters

- **obs** – ([float] or [int]) The current observation of the environment
- **state** – ([float]) The last states (used in recurrent policies)
- **mask** – ([float]) The last masks (used in recurrent policies)

Returns ([float]) actions

value (*obs, action, state=None, mask=None*)
Returns the value for a single step

Parameters

- **obs** – ([float] or [int]) The current observation of the environment
- **action** – ([float] or [int]) The taken action

- **state** – ([float]) The last states (used in recurrent policies)
- **mask** – ([float]) The last masks (used in recurrent policies)

Returns ([float]) The associated value of the action

class `stable_baselines.ddpg.CnnPolicy` (*sess, ob_space, ac_space, n_env, n_steps, n_batch, reuse=False, **kwargs*)

Policy object that implements actor critic, using a CNN (the nature CNN)

Parameters

- **sess** – (TensorFlow session) The current TensorFlow session
- **ob_space** – (Gym Space) The observation space of the environment
- **ac_space** – (Gym Space) The action space of the environment
- **n_env** – (int) The number of environments to run
- **n_steps** – (int) The number of steps to run for each environment
- **n_batch** – (int) The number of batch to run ($n_{\text{envs}} * n_{\text{steps}}$)
- **reuse** – (bool) If the policy is reusable or not
- **kwargs** – (dict) Extra keyword arguments for the nature CNN feature extraction

make_actor (*obs=None, reuse=False, scope='pi'*)
creates an actor object

Parameters

- **obs** – (TensorFlow Tensor) The observation placeholder (can be None for default placeholder)
- **reuse** – (bool) whether or not to reuse parameters
- **scope** – (str) the scope name of the actor

Returns (TensorFlow Tensor) the output tensor

make_critic (*obs=None, action=None, reuse=False, scope='qf'*)
creates a critic object

Parameters

- **obs** – (TensorFlow Tensor) The observation placeholder (can be None for default placeholder)
- **action** – (TensorFlow Tensor) The action placeholder (can be None for default placeholder)
- **reuse** – (bool) whether or not to reuse parameters
- **scope** – (str) the scope name of the critic

Returns (TensorFlow Tensor) the output tensor

proba_step (*obs, state=None, mask=None*)
Returns the action probability for a single step

Parameters

- **obs** – ([float] or [int]) The current observation of the environment
- **state** – ([float]) The last states (used in recurrent policies)
- **mask** – ([float]) The last masks (used in recurrent policies)

Returns ([float]) the action probability

step (*obs*, *state=None*, *mask=None*)

Returns the policy for a single step

Parameters

- **obs** – ([float] or [int]) The current observation of the environment
- **state** – ([float]) The last states (used in recurrent policies)
- **mask** – ([float]) The last masks (used in recurrent policies)

Returns ([float]) actions

value (*obs*, *action*, *state=None*, *mask=None*)

Returns the value for a single step

Parameters

- **obs** – ([float] or [int]) The current observation of the environment
- **action** – ([float] or [int]) The taken action
- **state** – ([float]) The last states (used in recurrent policies)
- **mask** – ([float]) The last masks (used in recurrent policies)

Returns ([float]) The associated value of the action

class `stable_baselines.ddpg.LnCnnPolicy` (*sess*, *ob_space*, *ac_space*, *n_env*, *n_steps*, *n_batch*, *reuse=False*, ***kwargs*)

Policy object that implements actor critic, using a CNN (the nature CNN), with layer normalisation

Parameters

- **sess** – (TensorFlow session) The current TensorFlow session
- **ob_space** – (Gym Space) The observation space of the environment
- **ac_space** – (Gym Space) The action space of the environment
- **n_env** – (int) The number of environments to run
- **n_steps** – (int) The number of steps to run for each environment
- **n_batch** – (int) The number of batch to run (*n_envs* * *n_steps*)
- **reuse** – (bool) If the policy is reusable or not
- **_kwargs** – (dict) Extra keyword arguments for the nature CNN feature extraction

make_actor (*obs=None*, *reuse=False*, *scope='pi'*)

creates an actor object

Parameters

- **obs** – (TensorFlow Tensor) The observation placeholder (can be None for default placeholder)
- **reuse** – (bool) whether or not to reuse parameters
- **scope** – (str) the scope name of the actor

Returns (TensorFlow Tensor) the output tensor

make_critic (*obs=None*, *action=None*, *reuse=False*, *scope='qf'*)

creates a critic object

Parameters

- **obs** – (TensorFlow Tensor) The observation placeholder (can be None for default placeholder)
- **action** – (TensorFlow Tensor) The action placeholder (can be None for default placeholder)
- **reuse** – (bool) whether or not to reuse parameters
- **scope** – (str) the scope name of the critic

Returns (TensorFlow Tensor) the output tensor

proba_step (*obs*, *state=None*, *mask=None*)

Returns the action probability for a single step

Parameters

- **obs** – ([float] or [int]) The current observation of the environment
- **state** – ([float]) The last states (used in recurrent policies)
- **mask** – ([float]) The last masks (used in recurrent policies)

Returns ([float]) the action probability

step (*obs*, *state=None*, *mask=None*)

Returns the policy for a single step

Parameters

- **obs** – ([float] or [int]) The current observation of the environment
- **state** – ([float]) The last states (used in recurrent policies)
- **mask** – ([float]) The last masks (used in recurrent policies)

Returns ([float]) actions

value (*obs*, *action*, *state=None*, *mask=None*)

Returns the value for a single step

Parameters

- **obs** – ([float] or [int]) The current observation of the environment
- **action** – ([float] or [int]) The taken action
- **state** – ([float]) The last states (used in recurrent policies)
- **mask** – ([float]) The last masks (used in recurrent policies)

Returns ([float]) The associated value of the action

1.14.6 Action and Parameters Noise

```
class stable_baselines.ddpg.AdaptiveParamNoiseSpec (initial_stddev=0.1,          de-
                                                    desired_action_stddev=0.1,    adop-
                                                    tion_coefficient=1.01)
```

Implements adaptive parameter noise

Parameters

- **initial_stddev** – (float) the initial value for the standard deviation of the noise
- **desired_action_stddev** – (float) the desired value for the standard deviation of the noise

- **adoption_coefficient** – (float) the update coefficient for the standard deviation of the noise

adapt (*distance*)

update the standard deviation for the parameter noise

Parameters **distance** – (float) the noise distance applied to the parameters

get_stats ()

return the standard deviation for the parameter noise

Returns (dict) the stats of the noise

class `stable_baselines.ddpg.NormalActionNoise` (*mean, sigma*)

A gaussian action noise

Parameters

- **mean** – (float) the mean value of the noise
- **sigma** – (float) the scale of the noise (std here)

reset ()

call end of episode reset for the noise

class `stable_baselines.ddpg.OrnsteinUhlenbeckActionNoise` (*mean, sigma, theta=0.15, dt=0.01, initial_noise=None*)

A Ornstein Uhlenbeck action noise, this is designed to aproximate brownian motion with friction.

Based on <http://math.stackexchange.com/questions/1287634/implementing-ornstein-uhlenbeck-in-matlab>

Parameters

- **mean** – (float) the mean of the noise
- **sigma** – (float) the scale of the noise
- **theta** – (float) the rate of mean reversion
- **dt** – (float) the timestep for the noise
- **initial_noise** – ([float]) the initial value for the noise output, (if None: 0)

reset ()

reset the Ornstein Uhlenbeck noise, to the initial position

1.14.7 Custom Policy Network

Similarly to the example given in the [examples](#) page. You can easily define a custom architecture for the policy network:

```
import gym

from stable_baselines.ddpg.policies import FeedForwardPolicy
from stable_baselines.common.vec_env import DummyVecEnv
from stable_baselines import DDPG

# Custom MLP policy of three layers of size 128 each
class CustomPolicy(FeedForwardPolicy):
    def __init__(self, *args, **kwargs):
        super(CustomPolicy, self).__init__(*args, **kwargs,
                                           layers=[128, 128, 128],
```

(continues on next page)

(continued from previous page)

```

layer_norm=False,
feature_extraction="mlp")

# Create and wrap the environment
env = gym.make('Pendulum-v0')
env = DummyVecEnv([lambda: env])

model = DDPG(CustomPolicy, env, verbose=1)
# Train the agent
model.learn(total_timesteps=100000)

```

1.15 DQN

Deep Q Network (DQN) and its extensions (Double-DQN, Dueling-DQN, Prioritized Experience Replay).

Warning: The DQN model does not support `stable_baselines.common.policies`, as a result it must use its own policy models (see [DQN Policies](#)).

Available Policies

<i>MlpPolicy</i>	Policy object that implements DQN policy, using a MLP (2 layers of 64)
<i>LnMlpPolicy</i>	Policy object that implements DQN policy, using a MLP (2 layers of 64), with layer normalisation
<i>CnnPolicy</i>	Policy object that implements DQN policy, using a CNN (the nature CNN)
<i>LnCnnPolicy</i>	Policy object that implements DQN policy, using a CNN (the nature CNN), with layer normalisation

1.15.1 Notes

- Original paper: <https://arxiv.org/abs/1312.5602>

1.15.2 Can I use?

- Recurrent policies:
- Multi processing:
- Gym spaces:

Space	Action	Observation
Discrete		
Box		
MultiDiscrete		
MultiBinary		

1.15.3 Example

```
import gym

from stable_baselines.common.vec_env import DummyVecEnv
from stable_baselines.deepq.policies import MlpPolicy
from stable_baselines import DQN

env = gym.make('CartPole-v1')
env = DummyVecEnv([lambda: env])

model = DQN(MlpPolicy, env, verbose=1)
model.learn(total_timesteps=25000)
model.save("deepq_cartpole")

del model # remove to demonstrate saving and loading

model = DQN.load("deepq_cartpole")

obs = env.reset()
while True:
    action, _states = model.predict(obs)
    obs, rewards, dones, info = env.step(action)
    env.render()
```

With Atari:

```
from stable_baselines.common.atari_wrappers import make_atari
from stable_baselines.deepq.policies import MlpPolicy, CnnPolicy
from stable_baselines import DQN

env = make_atari('BreakoutNoFrameskip-v4')

model = DQN(CnnPolicy, env, verbose=1)
model.learn(total_timesteps=25000)
model.save("deepq_breakout")

del model # remove to demonstrate saving and loading

DQN.load("deepq_breakout")

obs = env.reset()
while True:
    action, _states = model.predict(obs)
    obs, rewards, dones, info = env.step(action)
    env.render()
```

1.15.4 Parameters

```
class stable_baselines.deepq.DQN(policy, env, gamma=0.99, learning_rate=0.0005,
    buffer_size=50000, exploration_fraction=0.1, exploration_final_eps=0.02,
    train_freq=1, batch_size=32, checkpoint_freq=10000, checkpoint_path=None,
    learning_starts=1000, target_network_update_freq=500, prioritized_replay=False,
    prioritized_replay_alpha=0.6, prioritized_replay_beta0=0.4,
    prioritized_replay_beta_iters=None, prioritized_replay_eps=1e-06,
    param_noise=False, verbose=0, tensorboard_log=None, _init_setup_model=True)
```

The DQN model class. DQN paper: <https://arxiv.org/pdf/1312.5602.pdf>

Parameters

- **policy** – (DQNPolicy or str) The policy model to use (MlpPolicy, CnnPolicy, LnMlpPolicy, ...)
- **env** – (Gym environment or str) The environment to learn from (if registered in Gym, can be str)
- **gamma** – (float) discount factor
- **learning_rate** – (float) learning rate for adam optimizer
- **buffer_size** – (int) size of the replay buffer
- **exploration_fraction** – (float) fraction of entire training period over which the exploration rate is annealed
- **exploration_final_eps** – (float) final value of random action probability
- **train_freq** – (int) update the model every *train_freq* steps. set to None to disable printing
- **batch_size** – (int) size of a batched sampled from replay buffer for training
- **checkpoint_freq** – (int) how often to save the model. This is so that the best version is restored at the end of the training. If you do not wish to restore the best version at the end of the training set this variable to None.
- **checkpoint_path** – (str) replacement path used if you need to log to somewhere else than a temporary directory.
- **learning_starts** – (int) how many steps of the model to collect transitions for before learning starts
- **target_network_update_freq** – (int) update the target network every *target_network_update_freq* steps.
- **prioritized_replay** – (bool) if True prioritized replay buffer will be used.
- **prioritized_replay_alpha** – (float) alpha parameter for prioritized replay buffer
- **prioritized_replay_beta0** – (float) initial value of beta for prioritized replay buffer
- **prioritized_replay_beta_iters** – (int) number of iterations over which beta will be annealed from initial value to 1.0. If set to None equals to max_timesteps.
- **prioritized_replay_eps** – (float) epsilon to add to the TD errors when updating priorities.
- **param_noise** – (bool) Whether or not to apply noise to the parameters of the policy.

- **verbose** – (int) the verbosity level: 0 none, 1 training information, 2 tensorflow debug
- **tensorboard_log** – (str) the log location for tensorboard (if None, no logging)
- **_init_setup_model** – (bool) Whether or not to build the network at the creation of the instance

action_probability (*observation, state=None, mask=None*)

Get the model's action probability distribution from an observation

Parameters

- **observation** – (np.ndarray) the input observation
- **state** – (np.ndarray) The last states (can be None, used in recurrent policies)
- **mask** – (np.ndarray) The last masks (can be None, used in recurrent policies)

Returns (np.ndarray) the model's action probability distribution

get_env ()

returns the current environment (can be None if not defined)

Returns (Gym Environment) The current environment

learn (*total_timesteps, callback=None, seed=None, log_interval=100, tb_log_name='DQN'*)

Return a trained model.

Parameters

- **total_timesteps** – (int) The total number of samples to train on
- **seed** – (int) The initial seed for training, if None: keep current seed
- **callback** – (function (dict, dict)) function called at every steps with state of the algorithm. It takes the local and global variables.
- **log_interval** – (int) The number of timesteps before logging.
- **tb_log_name** – (str) the name of the run for tensorboard log

Returns (BaseRLModel) the trained model

classmethod load (*load_path, env=None, **kwargs*)

Load the model from file

Parameters

- **load_path** – (str) the saved parameter location
- **env** – (Gym Environment) the new environment to run the loaded model on (can be None if you only need prediction from a trained model)
- **kwargs** – extra arguments to change the model when loading

predict (*observation, state=None, mask=None, deterministic=True*)

Get the model's action from an observation

Parameters

- **observation** – (np.ndarray) the input observation
- **state** – (np.ndarray) The last states (can be None, used in recurrent policies)
- **mask** – (np.ndarray) The last masks (can be None, used in recurrent policies)
- **deterministic** – (bool) Whether or not to return deterministic actions.

Returns (np.ndarray, np.ndarray) the model's action and the next state (used in recurrent policies)

save (*save_path*)

Save the current parameters to file

Parameters **save_path** – (str) the save location

set_env (*env*)

Checks the validity of the environment, and if it is coherent, set it as the current environment.

Parameters **env** – (Gym Environment) The environment for learning a policy

setup_model ()

Create all the functions and tensorflow graphs necessary to train the model

1.15.5 DQN Policies

class `stable_baselines.deepq.MlpPolicy` (*sess, ob_space, ac_space, n_env, n_steps, n_batch, reuse=False, obs_phs=None, dueling=True, **kwargs*)

Policy object that implements DQN policy, using a MLP (2 layers of 64)

Parameters

- **sess** – (TensorFlow session) The current TensorFlow session
- **ob_space** – (Gym Space) The observation space of the environment
- **ac_space** – (Gym Space) The action space of the environment
- **n_env** – (int) The number of environments to run
- **n_steps** – (int) The number of steps to run for each environment
- **n_batch** – (int) The number of batch to run ($n_{\text{envs}} * n_{\text{steps}}$)
- **reuse** – (bool) If the policy is reusable or not
- **obs_phs** – (TensorFlow Tensor, TensorFlow Tensor) a tuple containing an override for observation placeholder and the processed observation placeholder respectively
- **dueling** – (bool) if true double the output MLP to compute a baseline for action scores
- **_kwargs** – (dict) Extra keyword arguments for the nature CNN feature extraction

proba_step (*obs, state=None, mask=None*)

Returns the action probability for a single step

Parameters

- **obs** – (np.ndarray float or int) The current observation of the environment
- **state** – (np.ndarray float) The last states (used in recurrent policies)
- **mask** – (np.ndarray float) The last masks (used in recurrent policies)

Returns (np.ndarray float) the action probability

step (*obs, state=None, mask=None, deterministic=True*)

Returns the q_values for a single step

Parameters

- **obs** – (np.ndarray float or int) The current observation of the environment

- **state** – (np.ndarray float) The last states (used in recurrent policies)
- **mask** – (np.ndarray float) The last masks (used in recurrent policies)
- **deterministic** – (bool) Whether or not to return deterministic actions.

Returns (np.ndarray int, np.ndarray float, np.ndarray float) actions, q_values, states

class stable_baselines.deepq.**LnMlpPolicy**(*sess, ob_space, ac_space, n_env, n_steps, n_batch, reuse=False, obs_phs=None, dueling=True, **kwargs*)

Policy object that implements DQN policy, using a MLP (2 layers of 64), with layer normalisation

Parameters

- **sess** – (TensorFlow session) The current TensorFlow session
- **ob_space** – (Gym Space) The observation space of the environment
- **ac_space** – (Gym Space) The action space of the environment
- **n_env** – (int) The number of environments to run
- **n_steps** – (int) The number of steps to run for each environment
- **n_batch** – (int) The number of batch to run ($n_{\text{envs}} * n_{\text{steps}}$)
- **reuse** – (bool) If the policy is reusable or not
- **obs_phs** – (TensorFlow Tensor, TensorFlow Tensor) a tuple containing an override for observation placeholder and the processed observation placeholder respectively
- **dueling** – (bool) if true double the output MLP to compute a baseline for action scores
- **_kwargs** – (dict) Extra keyword arguments for the nature CNN feature extraction

proba_step (*obs, state=None, mask=None*)

Returns the action probability for a single step

Parameters

- **obs** – (np.ndarray float or int) The current observation of the environment
- **state** – (np.ndarray float) The last states (used in recurrent policies)
- **mask** – (np.ndarray float) The last masks (used in recurrent policies)

Returns (np.ndarray float) the action probability

step (*obs, state=None, mask=None, deterministic=True*)

Returns the q_values for a single step

Parameters

- **obs** – (np.ndarray float or int) The current observation of the environment
- **state** – (np.ndarray float) The last states (used in recurrent policies)
- **mask** – (np.ndarray float) The last masks (used in recurrent policies)
- **deterministic** – (bool) Whether or not to return deterministic actions.

Returns (np.ndarray int, np.ndarray float, np.ndarray float) actions, q_values, states

class stable_baselines.deepq.**CnnPolicy**(*sess, ob_space, ac_space, n_env, n_steps, n_batch, reuse=False, obs_phs=None, dueling=True, **kwargs*)

Policy object that implements DQN policy, using a CNN (the nature CNN)

Parameters

- **sess** – (TensorFlow session) The current TensorFlow session
- **ob_space** – (Gym Space) The observation space of the environment
- **ac_space** – (Gym Space) The action space of the environment
- **n_env** – (int) The number of environments to run
- **n_steps** – (int) The number of steps to run for each environment
- **n_batch** – (int) The number of batch to run ($n_{\text{envs}} * n_{\text{steps}}$)
- **reuse** – (bool) If the policy is reusable or not
- **obs_phs** – (TensorFlow Tensor, TensorFlow Tensor) a tuple containing an override for observation placeholder and the processed observation placeholder respectively
- **dueling** – (bool) if true double the output MLP to compute a baseline for action scores
- **_kwargs** – (dict) Extra keyword arguments for the nature CNN feature extraction

proba_step (*obs*, *state=None*, *mask=None*)

Returns the action probability for a single step

Parameters

- **obs** – (np.ndarray float or int) The current observation of the environment
- **state** – (np.ndarray float) The last states (used in recurrent policies)
- **mask** – (np.ndarray float) The last masks (used in recurrent policies)

Returns (np.ndarray float) the action probability

step (*obs*, *state=None*, *mask=None*, *deterministic=True*)

Returns the q_values for a single step

Parameters

- **obs** – (np.ndarray float or int) The current observation of the environment
- **state** – (np.ndarray float) The last states (used in recurrent policies)
- **mask** – (np.ndarray float) The last masks (used in recurrent policies)
- **deterministic** – (bool) Whether or not to return deterministic actions.

Returns (np.ndarray int, np.ndarray float, np.ndarray float) actions, q_values, states

class stable_baselines.deepq.**LncnnPolicy** (*sess*, *ob_space*, *ac_space*, *n_env*, *n_steps*, *n_batch*, *reuse=False*, *obs_phs=None*, *dueling=True*, ***kwargs*)

Policy object that implements DQN policy, using a CNN (the nature CNN), with layer normalisation

Parameters

- **sess** – (TensorFlow session) The current TensorFlow session
- **ob_space** – (Gym Space) The observation space of the environment
- **ac_space** – (Gym Space) The action space of the environment
- **n_env** – (int) The number of environments to run
- **n_steps** – (int) The number of steps to run for each environment
- **n_batch** – (int) The number of batch to run ($n_{\text{envs}} * n_{\text{steps}}$)

- **reuse** – (bool) If the policy is reusable or not
- **obs_phs** – (TensorFlow Tensor, TensorFlow Tensor) a tuple containing an override for observation placeholder and the processed observation placeholder respectively
- **dueling** – (bool) if true double the output MLP to compute a baseline for action scores
- **_kwargs** – (dict) Extra keyword arguments for the nature CNN feature extraction

proba_step (*obs, state=None, mask=None*)

Returns the action probability for a single step

Parameters

- **obs** – (np.ndarray float or int) The current observation of the environment
- **state** – (np.ndarray float) The last states (used in recurrent policies)
- **mask** – (np.ndarray float) The last masks (used in recurrent policies)

Returns (np.ndarray float) the action probability

step (*obs, state=None, mask=None, deterministic=True*)

Returns the `q_values` for a single step

Parameters

- **obs** – (np.ndarray float or int) The current observation of the environment
- **state** – (np.ndarray float) The last states (used in recurrent policies)
- **mask** – (np.ndarray float) The last masks (used in recurrent policies)
- **deterministic** – (bool) Whether or not to return deterministic actions.

Returns (np.ndarray int, np.ndarray float, np.ndarray float) actions, `q_values`, states

1.15.6 Custom Policy Network

Similarly to the example given in the [examples](#) page. You can easily define a custom architecture for the policy network:

```
import gym

from stable_baselines.deepq.policies import FeedForwardPolicy
from stable_baselines.common.vec_env import DummyVecEnv
from stable_baselines import DQN

# Custom MLP policy of three layers of size 128 each
class CustomPolicy(FeedForwardPolicy):
    def __init__(self, *args, **kwargs):
        super(CustomPolicy, self).__init__(*args, **kwargs,
                                           layers=[128, 128, 128],
                                           layer_norm=False,
                                           feature_extraction="mlp")

# Create and wrap the environment
env = gym.make('LunarLander-v2')
env = DummyVecEnv([lambda: env])

model = DQN(CustomPolicy, env, verbose=1)
# Train the agent
model.learn(total_timesteps=100000)
```

1.16 GAIL

Generative Adversarial Imitation Learning (GAIL)

1.16.1 Notes

- Original paper: <https://arxiv.org/abs/1606.03476>

1.16.2 If you want to train an imitation learning agent

Step 1: Download expert data

Download the expert data into `./data`, [download link](#)

Step 2: Run GAIL

Run with single thread:

```
python -m stable_baselines.gail.run_mujoco
```

Run with multiple threads:

```
mpirun -np 16 python -m stable_baselines.gail.run_mujoco
```

See help (`-h`) for more options.

In case you want to run Behavior Cloning (BC)

```
python -m stable_baselines.gail.behavior_clone
```

See help (`-h`) for more options.

OpenAI Maintainers:

- Yuan-Hong Liao, [andrewliao11_at_gmail_dot_com](mailto:andrewliao11@gmail.com)
- Ryan Julian, ryanjulian_at_gmail_dot_com

Others

Thanks to the open source:

- [@openai/imitation](#)
- [@carpedm20/deep-rl-tensorflow](#)

1.16.3 Can I use?

- Recurrent policies:
- Multi processing: (using MPI)
- Gym spaces:

Space	Action	Observation
Discrete		
Box		
MultiDiscrete		
MultiBinary		

1.16.4 Parameters

```
class stable_baselines.gail.GAIL(policy, env, pretrained_weight=False, hid-
    den_size_adversary=100, adversary_entcoeff=0.001,
    expert_dataset=None, save_per_iter=1, check-
    point_dir='/tmp/gail/ckpt/', g_step=1, d_step=1,
    task_name='task_name', d_stepsize=0.0003, verbose=0,
    _init_setup_model=True, **kwargs)
```

Generative Adversarial Imitation Learning (GAIL)

Parameters

- **policy** – (ActorCriticPolicy or str) The policy model to use (MlpPolicy, CnnPolicy, CnnLstmPolicy, ...)
- **env** – (Gym environment or str) The environment to learn from (if registered in Gym, can be str)
- **gamma** – (float) the discount value
- **timesteps_per_batch** – (int) the number of timesteps to run per batch (horizon)
- **max_kl** – (float) the kullback leiber loss threshold
- **cg_iters** – (int) the number of iterations for the conjugate gradient calculation
- **lam** – (float) GAE factor
- **entcoeff** – (float) the weight for the entropy loss
- **cg_damping** – (float) the compute gradient dampening factor
- **vf_stepsize** – (float) the value function stepsize
- **vf_iters** – (int) the value function's number iterations for learning
- **pretrained_weight** – (str) the save location for the pretrained weights
- **hidden_size** – ([int]) the hidden dimension for the MLP
- **expert_dataset** – (Dset) the dataset manager
- **save_per_iter** – (int) the number of iterations before saving
- **checkpoint_dir** – (str) the location for saving checkpoints
- **g_step** – (int) number of steps to train policy in each epoch
- **d_step** – (int) number of steps to train discriminator in each epoch
- **task_name** – (str) the name of the task (can be None)
- **d_stepsize** – (float) the reward giver stepsize
- **verbose** – (int) the verbosity level: 0 none, 1 training information, 2 tensorflow debug
- **_init_setup_model** – (bool) Whether or not to build the network at the creation of the instance

action_probability (*observation, state=None, mask=None*)
 Get the model's action probability distribution from an observation

Parameters

- **observation** – (np.ndarray) the input observation
- **state** – (np.ndarray) The last states (can be None, used in recurrent policies)
- **mask** – (np.ndarray) The last masks (can be None, used in recurrent policies)

Returns (np.ndarray) the model's action probability distribution

get_env ()
 returns the current environment (can be None if not defined)

Returns (Gym Environment) The current environment

learn (*total_timesteps, callback=None, seed=None, log_interval=100, tb_log_name='GAIL'*)
 Return a trained model.

Parameters

- **total_timesteps** – (int) The total number of samples to train on
- **seed** – (int) The initial seed for training, if None: keep current seed
- **callback** – (function (dict, dict)) function called at every steps with state of the algorithm. It takes the local and global variables.
- **log_interval** – (int) The number of timesteps before logging.
- **tb_log_name** – (str) the name of the run for tensorboard log

Returns (BaseRLModel) the trained model

classmethod load (*load_path, env=None, **kwargs*)
 Load the model from file

Parameters

- **load_path** – (str) the saved parameter location
- **env** – (Gym Environment) the new environment to run the loaded model on (can be None if you only need prediction from a trained model)
- **kwargs** – extra arguments to change the model when loading

predict (*observation, state=None, mask=None, deterministic=False*)
 Get the model's action from an observation

Parameters

- **observation** – (np.ndarray) the input observation
- **state** – (np.ndarray) The last states (can be None, used in recurrent policies)
- **mask** – (np.ndarray) The last masks (can be None, used in recurrent policies)
- **deterministic** – (bool) Whether or not to return deterministic actions.

Returns (np.ndarray, np.ndarray) the model's action and the next state (used in recurrent policies)

save (*save_path*)
 Save the current parameters to file

Parameters **save_path** – (str) the save location

set_env (*env*)

Checks the validity of the environment, and if it is coherent, set it as the current environment.

Parameters *env* – (Gym Environment) The environment for learning a policy

setup_model ()

Create all the functions and tensorflow graphs necessary to train the model

1.17 HER

Hindsight Experience Replay (HER)

Warning: HER is not refactored yet. We are looking for contributors to help us.

1.17.1 How to use Hindsight Experience Replay

Getting started

Training an agent is very simple:

```
python -m stable_baselines.her.experiment.train
```

This will train a DDPG+HER agent on the `FetchReach` environment. You should see the success rate go up quickly to 1.0, which means that the agent achieves the desired goal in 100% of the cases. The training script logs other diagnostics as well and pickles the best policy so far (w.r.t. to its test success rate), the latest policy, and, if enabled, a history of policies every *K* epochs.

To inspect what the agent has learned, use the play script:

```
python -m stable_baselines.her.experiment.play /path/to/an/experiment/policy_best.pkl
```

You can try it right now with the results of the training step (the script prints out the path for you). This should visualize the current policy for 10 episodes and will also print statistics.

Reproducing results

In order to reproduce the results from [Plappert et al. \(2018\)](#), run the following command:

```
python -m stable_baselines.her.experiment.train --num_cpu 19
```

This will require a machine with sufficient amount of physical CPU cores. In our experiments, we used [Azure's D15v2 instances](#), which have 20 physical cores. We only scheduled the experiment on 19 of those to leave some head-room on the system.

1.17.2 Parameters

class `stable_baselines.her.HER` (*policy*, *env*, *verbose=0*, *_init_setup_model=True*)

action_probability (*observation*, *state=None*, *mask=None*)

Get the model's action probability distribution from an observation

Parameters

- **observation** – (np.ndarray) the input observation
- **state** – (np.ndarray) The last states (can be None, used in recurrent policies)
- **mask** – (np.ndarray) The last masks (can be None, used in recurrent policies)

Returns (np.ndarray) the model’s action probability distribution

learn (*total_timesteps, callback=None, seed=None, log_interval=100, tb_log_name='HER'*)
Return a trained model.

Parameters

- **total_timesteps** – (int) The total number of samples to train on
- **seed** – (int) The initial seed for training, if None: keep current seed
- **callback** – (function (dict, dict)) function called at every steps with state of the algorithm. It takes the local and global variables.
- **log_interval** – (int) The number of timesteps before logging.
- **tb_log_name** – (str) the name of the run for tensorboard log

Returns (BaseRLModel) the trained model

classmethod load (*load_path, env=None, **kwargs*)
Load the model from file

Parameters

- **load_path** – (str) the saved parameter location
- **env** – (Gym Environment) the new environment to run the loaded model on (can be None if you only need prediction from a trained model)
- **kwargs** – extra arguments to change the model when loading

predict (*observation, state=None, mask=None, deterministic=False*)
Get the model’s action from an observation

Parameters

- **observation** – (np.ndarray) the input observation
- **state** – (np.ndarray) The last states (can be None, used in recurrent policies)
- **mask** – (np.ndarray) The last masks (can be None, used in recurrent policies)
- **deterministic** – (bool) Whether or not to return deterministic actions.

Returns (np.ndarray, np.ndarray) the model’s action and the next state (used in recurrent policies)

save (*save_path*)
Save the current parameters to file

Parameters **save_path** – (str) the save location

setup_model ()
Create all the functions and tensorflow graphs necessary to train the model

1.18 PPO1

The [Proximal Policy Optimization](#) algorithm combines ideas from A2C (having multiple workers) and TRPO (it uses a trust region to improve the actor).

The main idea is that after an update, the new policy should be not too far from the *old* policy. For that, ppo uses clipping to avoid too large update.

Note: PPO2 is the implementation of OpenAI made for GPU. For multiprocessing, it uses vectorized environments compared to PPO1 which uses MPI.

1.18.1 Notes

- Original paper: <https://arxiv.org/abs/1707.06347>
- Clear explanation of PPO on Arxiv Insights channel: <https://www.youtube.com/watch?v=5P7I-xPq8u8>
- OpenAI blog post: <https://blog.openai.com/openai-baselines-ppo/>
- `mpirun -np 8 python -m stable_baselines.ppo1.run_atari` runs the algorithm for 40M frames = 10M timesteps on an Atari game. See help (-h) for more options.
- `python -m stable_baselines.ppo1.run_mujoco` runs the algorithm for 1M frames on a Mujoco environment.
- Train mujoco 3d humanoid (with optimal-ish hyperparameters): `mpirun -np 16 python -m stable_baselines.ppo1.run_humanoid --model-path=/path/to/model`
- Render the 3d humanoid: `python -m stable_baselines.ppo1.run_humanoid --play --model-path=/path/to/model`

1.18.2 Can I use?

- Recurrent policies:
- Multi processing: (using MPI)
- Gym spaces:

Space	Action	Observation
Discrete		
Box		
MultiDiscrete		
MultiBinary		

1.18.3 Example

```
import gym

from stable_baselines.common.policies import MlpPolicy, MlpLstmPolicy, MlpLnLstmPolicy
from stable_baselines.common.vec_env import DummyVecEnv
from stable_baselines import PPO1
```

(continues on next page)

(continued from previous page)

```

env = gym.make('CartPole-v1')
env = DummyVecEnv([lambda: env])

model = PPO1(MlpPolicy, env, verbose=1)
model.learn(total_timesteps=25000)
model.save("ppo1_cartpole")

del model # remove to demonstrate saving and loading

model = PPO1.load("ppo1_cartpole")

obs = env.reset()
while True:
    action, _states = model.predict(obs)
    obs, rewards, dones, info = env.step(action)
    env.render()

```

1.18.4 Parameters

class `stable_baselines.ppo1.PPO1` (*policy*, *env*, *gamma*=0.99, *timesteps_per_actorbatch*=256, *clip_param*=0.2, *entcoeff*=0.01, *optim_epochs*=4, *optim_stepsize*=0.001, *optim_batchsize*=64, *lam*=0.95, *adam_epsilon*=1e-05, *schedule*='linear', *verbose*=0, *tensorboard_log*=None, *_init_setup_model*=True)

Proximal Policy Optimization algorithm (MPI version). Paper: <https://arxiv.org/abs/1707.06347>

Parameters

- **env** – (Gym environment or str) The environment to learn from (if registered in Gym, can be str)
- **policy** – (ActorCriticPolicy or str) The policy model to use (MlpPolicy, CnnPolicy, CnnLstmPolicy, ...)
- **timesteps_per_actorbatch** – (int) timesteps per actor per update
- **clip_param** – (float) clipping parameter epsilon
- **entcoeff** – (float) the entropy loss weight
- **optim_epochs** – (float) the optimizer's number of epochs
- **optim_stepsize** – (float) the optimizer's stepsize
- **optim_batchsize** – (int) the optimizer's the batch size
- **gamma** – (float) discount factor
- **lam** – (float) advantage estimation
- **adam_epsilon** – (float) the epsilon value for the adam optimizer
- **schedule** – (str) The type of scheduler for the learning rate update ('linear', 'constant', 'double_linear_con', 'middle_drop' or 'double_middle_drop')
- **verbose** – (int) the verbosity level: 0 none, 1 training information, 2 tensorflow debug
- **tensorboard_log** – (str) the log location for tensorboard (if None, no logging)

- **_init_setup_model** – (bool) Whether or not to build the network at the creation of the instance

action_probability (*observation, state=None, mask=None*)

Get the model's action probability distribution from an observation

Parameters

- **observation** – (np.ndarray) the input observation
- **state** – (np.ndarray) The last states (can be None, used in recurrent policies)
- **mask** – (np.ndarray) The last masks (can be None, used in recurrent policies)

Returns (np.ndarray) the model's action probability distribution

get_env ()

returns the current environment (can be None if not defined)

Returns (Gym Environment) The current environment

learn (*total_timesteps, callback=None, seed=None, log_interval=100, tb_log_name='PPO1'*)

Return a trained model.

Parameters

- **total_timesteps** – (int) The total number of samples to train on
- **seed** – (int) The initial seed for training, if None: keep current seed
- **callback** – (function (dict, dict)) function called at every steps with state of the algorithm. It takes the local and global variables.
- **log_interval** – (int) The number of timesteps before logging.
- **tb_log_name** – (str) the name of the run for tensorboard log

Returns (BaseRLModel) the trained model

classmethod load (*load_path, env=None, **kwargs*)

Load the model from file

Parameters

- **load_path** – (str) the saved parameter location
- **env** – (Gym Environment) the new environment to run the loaded model on (can be None if you only need prediction from a trained model)
- **kwargs** – extra arguments to change the model when loading

predict (*observation, state=None, mask=None, deterministic=False*)

Get the model's action from an observation

Parameters

- **observation** – (np.ndarray) the input observation
- **state** – (np.ndarray) The last states (can be None, used in recurrent policies)
- **mask** – (np.ndarray) The last masks (can be None, used in recurrent policies)
- **deterministic** – (bool) Whether or not to return deterministic actions.

Returns (np.ndarray, np.ndarray) the model's action and the next state (used in recurrent policies)

save (*save_path*)

Save the current parameters to file

Parameters **save_path** – (str) the save location

set_env (*env*)

Checks the validity of the environment, and if it is coherent, set it as the current environment.

Parameters **env** – (Gym Environment) The environment for learning a policy

setup_model ()

Create all the functions and tensorflow graphs necessary to train the model

1.19 PPO2

The [Proximal Policy Optimization](#) algorithm combines ideas from A2C (having multiple workers) and TRPO (it uses a trust region to improve the actor).

The main idea is that after an update, the new policy should be not too far from the *old* policy. For that, ppo uses clipping to avoid too large update.

Note: PPO2 is the implementation of OpenAI made for GPU. For multiprocessing, it uses vectorized environments compared to PPO1 which uses MPI.

Note: PPO2 contains several modifications from the original algorithm not documented by OpenAI: value function is also clipped and advantages are normalized.

1.19.1 Notes

- Original paper: <https://arxiv.org/abs/1707.06347>
- Clear explanation of PPO on Arxiv Insights channel: <https://www.youtube.com/watch?v=5P7I-xPq8u8>
- OpenAI blog post: <https://blog.openai.com/openai-baselines-ppo/>
- **python -m stable_baselines.ppo2.run_atari** runs the algorithm for **40M** frames = 10M timesteps on an Atari game. See help (-h) for more options.
- **python -m stable_baselines.ppo2.run_mujoco** runs the algorithm for **1M** frames on a Mujoco environment.

1.19.2 Can I use?

- Recurrent policies:
- Multi processing:
- Gym spaces:

Space	Action	Observation
Discrete		
Box		
MultiDiscrete		
MultiBinary		

1.19.3 Example

Train a PPO agent on *CartPole-v1* using 4 processes.

```
import gym

from stable_baselines.common.policies import MlpPolicy
from stable_baselines.common.vec_env import SubprocVecEnv
from stable_baselines import PPO2

# multiprocessing environment
n_cpu = 4
env = SubprocVecEnv([lambda: gym.make('CartPole-v1') for i in range(n_cpu)])

model = PPO2(MlpPolicy, env, verbose=1)
model.learn(total_timesteps=25000)
model.save("ppo2_cartpole")

del model # remove to demonstrate saving and loading

model = PPO2.load("ppo2_cartpole")

# Enjoy trained agent
obs = env.reset()
while True:
    action, _states = model.predict(obs)
    obs, rewards, dones, info = env.step(action)
    env.render()
```

1.19.4 Parameters

```
class stable_baselines.ppo2.PPO2(policy, env, gamma=0.99, n_steps=128, ent_coef=0.01,
    learning_rate=0.00025, vf_coef=0.5, max_grad_norm=0.5,
    lam=0.95, nminibatches=4, noptepochs=4,
    cliprange=0.2, verbose=0, tensorboard_log=None,
    _init_setup_model=True)
```

Proximal Policy Optimization algorithm (GPU version). Paper: <https://arxiv.org/abs/1707.06347>

Parameters

- **policy** – (ActorCriticPolicy or str) The policy model to use (MlpPolicy, CnnPolicy, CnnLstmPolicy, ...)
- **env** – (Gym environment or str) The environment to learn from (if registered in Gym, can be str)
- **gamma** – (float) Discount factor
- **n_steps** – (int) The number of steps to run for each environment per update (i.e. batch size is $n_steps * n_env$ where n_env is number of environment copies running in parallel)

- **ent_coef** – (float) Entropy coefficient for the loss calculation
- **learning_rate** – (float or callable) The learning rate, it can be a function
- **vf_coef** – (float) Value function coefficient for the loss calculation
- **max_grad_norm** – (float) The maximum value for the gradient clipping
- **lam** – (float) Factor for trade-off of bias vs variance for Generalized Advantage Estimator
- **nminibatches** – (int) Number of training minibatches per update. For recurrent policies, the number of environments run in parallel should be a multiple of nminibatches.
- **noptepochs** – (int) Number of epoch when optimizing the surrogate
- **cliprange** – (float or callable) Clipping parameter, it can be a function
- **verbose** – (int) the verbosity level: 0 none, 1 training information, 2 tensorflow debug
- **tensorboard_log** – (str) the log location for tensorboard (if None, no logging)
- **_init_setup_model** – (bool) Whether or not to build the network at the creation of the instance

action_probability (*observation, state=None, mask=None*)

Get the model's action probability distribution from an observation

Parameters

- **observation** – (np.ndarray) the input observation
- **state** – (np.ndarray) The last states (can be None, used in recurrent policies)
- **mask** – (np.ndarray) The last masks (can be None, used in recurrent policies)

Returns (np.ndarray) the model's action probability distribution

get_env ()

returns the current environment (can be None if not defined)

Returns (Gym Environment) The current environment

learn (*total_timesteps, callback=None, seed=None, log_interval=1, tb_log_name='PPO2'*)

Return a trained model.

Parameters

- **total_timesteps** – (int) The total number of samples to train on
- **seed** – (int) The initial seed for training, if None: keep current seed
- **callback** – (function (dict, dict)) function called at every steps with state of the algorithm. It takes the local and global variables.
- **log_interval** – (int) The number of timesteps before logging.
- **tb_log_name** – (str) the name of the run for tensorboard log

Returns (BaseRLModel) the trained model

classmethod load (*load_path, env=None, **kwargs*)

Load the model from file

Parameters

- **load_path** – (str) the saved parameter location
- **env** – (Gym Environment) the new environment to run the loaded model on (can be None if you only need prediction from a trained model)

- **kwargs** – extra arguments to change the model when loading

predict (*observation*, *state=None*, *mask=None*, *deterministic=False*)

Get the model's action from an observation

Parameters

- **observation** – (np.ndarray) the input observation
- **state** – (np.ndarray) The last states (can be None, used in recurrent policies)
- **mask** – (np.ndarray) The last masks (can be None, used in recurrent policies)
- **deterministic** – (bool) Whether or not to return deterministic actions.

Returns (np.ndarray, np.ndarray) the model's action and the next state (used in recurrent policies)

save (*save_path*)

Save the current parameters to file

Parameters **save_path** – (str) the save location

set_env (*env*)

Checks the validity of the environment, and if it is coherent, set it as the current environment.

Parameters **env** – (Gym Environment) The environment for learning a policy

setup_model ()

Create all the functions and tensorflow graphs necessary to train the model

1.20 TRPO

Trust Region Policy Optimization (TRPO) is an iterative approach for optimizing policies with guaranteed monotonic improvement.

1.20.1 Notes

- Original paper: <https://arxiv.org/abs/1502.05477>
- OpenAI blog post: <https://blog.openai.com/openai-baselines-ppo/>
- `mpirun -np 16 python -m stable_baselines.trpo_mpi.run_atari` runs the algorithm for 40M frames = 10M timesteps on an Atari game. See help (-h) for more options.
- `python -m stable_baselines.trpo_mpi.run_mujoco` runs the algorithm for 1M timesteps on a Mujoco environment.

1.20.2 Can I use?

- Recurrent policies:
- Multi processing: (using MPI)
- Gym spaces:

Space	Action	Observation
Discrete		
Box		
MultiDiscrete		
MultiBinary		

1.20.3 Example

```
import gym

from stable_baselines.common.policies import MlpPolicy, MlpLstmPolicy, MlpLnLstmPolicy
from stable_baselines.common.vec_env import DummyVecEnv
from stable_baselines import TRPO

env = gym.make('CartPole-v1')
env = DummyVecEnv([lambda: env])

model = TRPO(MlpPolicy, env, verbose=1)
model.learn(total_timesteps=25000)
model.save("trpo_cartpole")

del model # remove to demonstrate saving and loading

model = TRPO.load("trpo_cartpole")

obs = env.reset()
while True:
    action, _states = model.predict(obs)
    obs, rewards, dones, info = env.step(action)
    env.render()
```

1.20.4 Parameters

```
class stable_baselines.trpo_mpi.TRPO(policy, env, gamma=0.99, timesteps_per_batch=1024,
                                     max_kl=0.01, cg_iters=10, lam=0.98, entco-
                                     eff=0.0, cg_damping=0.01, vf_stepsize=0.0003,
                                     vf_iters=3, verbose=0, tensorboard_log=None,
                                     _init_setup_model=True)
```

action_probability (observation, state=None, mask=None)
Get the model's action probability distribution from an observation

Parameters

- **observation** – (np.ndarray) the input observation
- **state** – (np.ndarray) The last states (can be None, used in recurrent policies)
- **mask** – (np.ndarray) The last masks (can be None, used in recurrent policies)

Returns (np.ndarray) the model's action probability distribution

get_env ()
returns the current environment (can be None if not defined)

Returns (Gym Environment) The current environment

learn (*total_timesteps*, *callback=None*, *seed=None*, *log_interval=100*, *tb_log_name='TRPO'*)

Return a trained model.

Parameters

- **total_timesteps** – (int) The total number of samples to train on
- **seed** – (int) The initial seed for training, if None: keep current seed
- **callback** – (function (dict, dict)) function called at every steps with state of the algorithm. It takes the local and global variables.
- **log_interval** – (int) The number of timesteps before logging.
- **tb_log_name** – (str) the name of the run for tensorboard log

Returns (BaseRLModel) the trained model

classmethod load (*load_path*, *env=None*, ***kwargs*)

Load the model from file

Parameters

- **load_path** – (str) the saved parameter location
- **env** – (Gym Environment) the new environment to run the loaded model on (can be None if you only need prediction from a trained model)
- **kwargs** – extra arguments to change the model when loading

predict (*observation*, *state=None*, *mask=None*, *deterministic=False*)

Get the model's action from an observation

Parameters

- **observation** – (np.ndarray) the input observation
- **state** – (np.ndarray) The last states (can be None, used in recurrent policies)
- **mask** – (np.ndarray) The last masks (can be None, used in recurrent policies)
- **deterministic** – (bool) Whether or not to return deterministic actions.

Returns (np.ndarray, np.ndarray) the model's action and the next state (used in recurrent policies)

save (*save_path*)

Save the current parameters to file

Parameters **save_path** – (str) the save location

set_env (*env*)

Checks the validity of the environment, and if it is coherent, set it as the current environment.

Parameters **env** – (Gym Environment) The environment for learning a policy

setup_model ()

Create all the functions and tensorflow graphs necessary to train the model

1.21 Probability Distributions

Probability distributions used for the different action spaces:

- CategoricalProbabilityDistribution -> Discrete

- `DiagGaussianProbabilityDistribution` -> `Box` (continuous actions)
- `MultiCategoricalProbabilityDistribution` -> `MultiDiscrete`
- `BernoulliProbabilityDistribution` -> `MultiBinary`

The policy networks output parameters for the distributions (named *flat* in the methods). Actions are then sampled from those distributions.

For instance, in the case of discrete actions. The policy network outputs probability of taking each action. The `CategoricalProbabilityDistribution` allows to sample from it, computes the entropy, the negative log probability (`neglogp`) and backpropagate the gradient.

In the case of continuous actions, a Gaussian distribution is used. The policy network outputs mean and (log) std of the distribution (assumed to be a `DiagGaussianProbabilityDistribution`).

class `stable_baselines.common.distributions.BernoulliProbabilityDistribution` (*logits*)

entropy ()

Returns shannon's entropy of the probability

Returns (float) the entropy

flatparam ()

Return the direct probabilities

Returns ([float]) the probabilities

classmethod fromflat (*flat*)

Create an instance of this from new bernoulli input

Parameters flat – ([float]) the bernoulli input data

Returns (ProbabilityDistribution) the instance from the given bernoulli input data

kl (*other*)

Calculates the Kullback-Leiber divergence from the given probability distribution

Parameters other – ([float]) the distribution to compare with

Returns (float) the KL divergence of the two distributions

mode ()

Returns the probability

Returns (Tensorflow Tensor) the deterministic action

neglogp (*x*)

returns the of the negative log likelihood

Parameters x – (str) the labels of each index

Returns ([float]) The negative log likelihood of the distribution

sample ()

returns a sample from the probability distribution

Returns (Tensorflow Tensor) the stochastic action

class `stable_baselines.common.distributions.BernoulliProbabilityDistributionType` (*size*)

param_shape ()

returns the shape of the input parameters

Returns ([int]) the shape

proba_distribution_from_latent (*pi_latent_vector*, *vf_latent_vector*, *init_scale=1.0*, *init_bias=0.0*)
 returns the probability distribution from latent values

Parameters

- **pi_latent_vector** – ([float]) the latent pi values
- **vf_latent_vector** – ([float]) the latent vf values
- **init_scale** – (float) the initial scale of the distribution
- **init_bias** – (float) the initial bias of the distribution

Returns (ProbabilityDistribution) the instance of the ProbabilityDistribution associated

probability_distribution_class ()
 returns the ProbabilityDistribution class of this type

Returns (Type ProbabilityDistribution) the probability distribution class associated

sample_dtype ()
 returns the type of the sampling

Returns (type) the type

sample_shape ()
 returns the shape of the sampling

Returns ([int]) the shape

class stable_baselines.common.distributions.CategoricalProbabilityDistribution (*logits*)

entropy ()
 Returns shannon's entropy of the probability

Returns (float) the entropy

flatparam ()
 Return the direct probabilities

Returns ([float]) the probabilities

classmethod fromflat (*flat*)
 Create an instance of this from new logits values

Parameters flat – ([float]) the categorical logits input

Returns (ProbabilityDistribution) the instance from the given categorical input

kl (*other*)
 Calculates the Kullback-Leiber divergence from the given probability distribution

Parameters other – ([float]) the distribution to compare with

Returns (float) the KL divergence of the two distributions

mode ()
 Returns the probability

Returns (Tensorflow Tensor) the deterministic action

neglogp (*x*)
 returns the of the negative log likelihood

Parameters x – (str) the labels of each index

Returns ([float]) The negative log likelihood of the distribution

sample()

returns a sample from the probability distribution

Returns (Tensorflow Tensor) the stochastic action

class `stable_baselines.common.distributions.CategoricalProbabilityDistributionType` (*n_cat*)

param_shape()

returns the shape of the input parameters

Returns ([int]) the shape

proba_distribution_from_latent (*pi_latent_vector*, *vf_latent_vector*, *init_scale=1.0*,
init_bias=0.0)

returns the probability distribution from latent values

Parameters

- **pi_latent_vector** – ([float]) the latent pi values
- **vf_latent_vector** – ([float]) the latent vf values
- **init_scale** – (float) the initial scale of the distribution
- **init_bias** – (float) the initial bias of the distribution

Returns (ProbabilityDistribution) the instance of the ProbabilityDistribution associated

probability_distribution_class()

returns the ProbabilityDistribution class of this type

Returns (Type ProbabilityDistribution) the probability distribution class associated

sample_dtype()

returns the type of the sampling

Returns (type) the type

sample_shape()

returns the shape of the sampling

Returns ([int]) the shape

class `stable_baselines.common.distributions.DiagGaussianProbabilityDistribution` (*flat*)

entropy()

Returns shannon's entropy of the probability

Returns (float) the entropy

flatparam()

Return the direct probabilities

Returns ([float]) the probabilities

classmethod fromflat (*flat*)

Create an instance of this from new multivariate gaussian input

Parameters **flat** – ([float]) the multivariate gaussian input data

Returns (ProbabilityDistribution) the instance from the given multivariate gaussian input data

kl (*other*)

Calculates the Kullback-Leiber divergence from the given probability distribution

Parameters *other* – ([float]) the distribution to compare with

Returns (float) the KL divergence of the two distributions

mode ()
Returns the probability

Returns (Tensorflow Tensor) the deterministic action

neglogp (*x*)
returns the of the negative log likelihood

Parameters *x* – (str) the labels of each index

Returns ([float]) The negative log likelihood of the distribution

sample ()
returns a sample from the probabilty distribution

Returns (Tensorflow Tensor) the stochastic action

class `stable_baselines.common.distributions.DiagGaussianProbabilityDistributionType` (*size*)

param_shape ()
returns the shape of the input parameters

Returns ([int]) the shape

proba_distribution_from_flat (*flat*)
returns the probability distribution from flat probabilities

Parameters *flat* – ([float]) the flat probabilities

Returns (ProbabilityDistribution) the instance of the ProbabilityDistribution associated

proba_distribution_from_latent (*pi_latent_vector*, *vf_latent_vector*, *init_scale=1.0*, *init_bias=0.0*)
returns the probability distribution from latent values

Parameters

- **pi_latent_vector** – ([float]) the latent pi values
- **vf_latent_vector** – ([float]) the latent vf values
- **init_scale** – (float) the initial scale of the distribution
- **init_bias** – (float) the initial bias of the distribution

Returns (ProbabilityDistribution) the instance of the ProbabilityDistribution associated

probability_distribution_class ()
returns the ProbabilityDistribution class of this type

Returns (Type ProbabilityDistribution) the probability distribution class associated

sample_dtype ()
returns the type of the sampling

Returns (type) the type

sample_shape ()
returns the shape of the sampling

Returns ([int]) the shape

```

class stable_baselines.common.distributions.MultiCategoricalProbabilityDistribution(nvec,
                                                                              flat)

    entropy()
        Returns shannon's entropy of the probability
        Returns (float) the entropy

    flatparam()
        Return the direct probabilities
        Returns ([float]) the probabilities

    classmethod fromflat(flat)
        Create an instance of this from new logits values
        Parameters flat – ([float]) the multi categorical logits input
        Returns (ProbabilityDistribution) the instance from the given multi categorical input

    kl(other)
        Calculates the Kullback-Leiber divergence from the given probability distribution
        Parameters other – ([float]) the distribution to compare with
        Returns (float) the KL divergence of the two distributions

    mode()
        Returns the probability
        Returns (Tensorflow Tensor) the deterministic action

    neglogp(x)
        returns the of the negative log likelihood
        Parameters x – (str) the labels of each index
        Returns ([float]) The negative log likelihood of the distribution

    sample()
        returns a sample from the probability distribution
        Returns (Tensorflow Tensor) the stochastic action

class stable_baselines.common.distributions.MultiCategoricalProbabilityDistributionType(n_v,
                                                                              n_a)

    param_shape()
        returns the shape of the input parameters
        Returns ([int]) the shape

    proba_distribution_from_flat(flat)
        Returns the probability distribution from flat probabilities flat: flattened vector of parameters of probability
        distribution
        Parameters flat – ([float]) the flat probabilities
        Returns (ProbabilityDistribution) the instance of the ProbabilityDistribution associated

    proba_distribution_from_latent(pi_latent_vector, vf_latent_vector, init_scale=1.0,
                                init_bias=0.0)
        returns the probability distribution from latent values
        Parameters
        • pi_latent_vector – ([float]) the latent pi values

```

- **vf_latent_vector** – ([float]) the latent vf values
- **init_scale** – (float) the initial scale of the distribution
- **init_bias** – (float) the initial bias of the distribution

Returns (ProbabilityDistribution) the instance of the ProbabilityDistribution associated

probability_distribution_class ()

returns the ProbabilityDistribution class of this type

Returns (Type ProbabilityDistribution) the probability distribution class associated

sample_dtype ()

returns the type of the sampling

Returns (type) the type

sample_shape ()

returns the shape of the sampling

Returns ([int]) the shape

class `stable_baselines.common.distributions.ProbabilityDistribution`

A particular probability distribution

entropy ()

Returns shannon's entropy of the probability

Returns (float) the entropy

flatparam ()

Return the direct probabilities

Returns ([float]) the probabilities

kl (*other*)

Calculates the Kullback-Leiber divergence from the given probability distribution

Parameters *other* – ([float]) the distribution to compare with

Returns (float) the KL divergence of the two distributions

logp (*x*)

returns the of the log likelihood

Parameters *x* – (str) the labels of each index

Returns ([float]) The log likelihood of the distribution

mode ()

Returns the probability

Returns (Tensorflow Tensor) the deterministic action

neglogp (*x*)

returns the of the negative log likelihood

Parameters *x* – (str) the labels of each index

Returns ([float]) The negative log likelihood of the distribution

sample ()

returns a sample from the probability distribution

Returns (Tensorflow Tensor) the stochastic action

class `stable_baselines.common.distributions.ProbabilityDistributionType`
Parametrized family of probability distributions

param_placeholder (*prepend_shape*, *name=None*)
returns the TensorFlow placeholder for the input parameters

Parameters

- **prepend_shape** – ([int]) the prepend shape
- **name** – (str) the placeholder name

Returns (TensorFlow Tensor) the placeholder

param_shape ()
returns the shape of the input parameters

Returns ([int]) the shape

proba_distribution_from_flat (*flat*)
Returns the probability distribution from flat probabilities flat: flattened vector of parameters of probability distribution

Parameters **flat** – ([float]) the flat probabilities

Returns (ProbabilityDistribution) the instance of the ProbabilityDistribution associated

proba_distribution_from_latent (*pi_latent_vector*, *vf_latent_vector*, *init_scale=1.0*, *init_bias=0.0*)
returns the probability distribution from latent values

Parameters

- **pi_latent_vector** – ([float]) the latent pi values
- **vf_latent_vector** – ([float]) the latent vf values
- **init_scale** – (float) the initial scale of the distribution
- **init_bias** – (float) the initial bias of the distribution

Returns (ProbabilityDistribution) the instance of the ProbabilityDistribution associated

probability_distribution_class ()
returns the ProbabilityDistribution class of this type

Returns (Type ProbabilityDistribution) the probability distribution class associated

sample_dtype ()
returns the type of the sampling

Returns (type) the type

sample_placeholder (*prepend_shape*, *name=None*)
returns the TensorFlow placeholder for the sampling

Parameters

- **prepend_shape** – ([int]) the prepend shape
- **name** – (str) the placeholder name

Returns (TensorFlow Tensor) the placeholder

sample_shape ()
returns the shape of the sampling

Returns ([int]) the shape

`stable_baselines.common.distributions.make_proba_dist_type(ac_space)`
 return an instance of ProbabilityDistributionType for the correct type of action space

Parameters `ac_space` – (Gym Space) the input action space

Returns (ProbabilityDistributionType) the appropriate instance of a ProbabilityDistributionType

`stable_baselines.common.distributions.shape_el(tensor, index)`
 get the shape of a TensorFlow Tensor element

Parameters

- **tensor** – (TensorFlow Tensor) the input tensor
- **index** – (int) the element

Returns ([int]) the shape

1.22 Tensorflow Utils

`stable_baselines.common.tf_util.conv2d(input_tensor, num_filters, name, filter_size=(3, 3), stride=(1, 1), pad='SAME', dtype=<MagicMock id='140270955926640'>, collections=None, summary_tag=None)`

Creates a 2d convolutional layer for TensorFlow

Parameters

- **input_tensor** – (TensorFlow Tensor) The input tensor for the convolution
- **num_filters** – (int) The number of filters
- **name** – (str) The TensorFlow variable scope
- **filter_size** – (tuple) The filter size
- **stride** – (tuple) The stride of the convolution
- **pad** – (str) The padding type ('VALID' or 'SAME')
- **dtype** – (type) The data type for the Tensors
- **collections** – (list) List of graph collections keys to add the Variable to
- **summary_tag** – (str) image summary name, can be None for no image summary

Returns (TensorFlow Tensor) 2d convolutional layer

`stable_baselines.common.tf_util.display_var_info(_vars)`
 log variable information, for debug purposes

Parameters `_vars` – ([TensorFlow Tensor]) the variables

`stable_baselines.common.tf_util.flatgrad(loss, var_list, clip_norm=None)`
 calculates the gradient and flattens it

Parameters

- **loss** – (float) the loss value
- **var_list** – ([TensorFlow Tensor]) the variables
- **clip_norm** – (float) clip the gradients (disabled if None)

Returns ([TensorFlow Tensor]) flattend gradient

`stable_baselines.common.tf_util.flattenallbut0 (tensor)`

flatten all the dimension, except from the first one

Parameters **tensor** – (TensorFlow Tensor) the input tensor

Returns (TensorFlow Tensor) the flattened tensor

`stable_baselines.common.tf_util.function (inputs, outputs, updates=None, givens=None)`

Just like Theano function. Take a bunch of tensorflow placeholders and expressions computed based on those placeholders and produces $f(\text{inputs}) \rightarrow \text{outputs}$. Function f takes values to be fed to the input's placeholders and produces the values of the expressions in outputs.

Input values can be passed in the same order as inputs or can be provided as kwargs based on placeholder name (passed to constructor or accessible via `placeholder.op.name`).

Example: `x = tf.placeholder(tf.int32, (), name="x") y = tf.placeholder(tf.int32, (), name="y") z = 3 * x + 2 * y`
`lin = function([x, y], z, givens={y: 0})`

with single_threaded_session(): initialize()

`assert lin(2) == 6 assert lin(x=3) == 9 assert lin(2, 2) == 10 assert lin(x=2, y=3) == 12`

Parameters

- **inputs** – (TensorFlow Tensor or Object with `make_feed_dict`) list of input arguments
- **outputs** – (TensorFlow Tensor) list of outputs or a single output to be returned from function. Returned value will also have the same shape.
- **updates** – (list) update functions
- **givens** – (dict) the values known for the output

`stable_baselines.common.tf_util.get_available_gpus ()`

Return a list of all the available GPUs

Returns ([str]) the GPUs available

`stable_baselines.common.tf_util.get_globals_vars (name)`

returns the trainable variables

Parameters **name** – (str) the scope

Returns ([TensorFlow Variable])

`stable_baselines.common.tf_util.get_trainable_vars (name)`

returns the trainable variables

Parameters **name** – (str) the scope

Returns ([TensorFlow Variable])

`stable_baselines.common.tf_util.huber_loss (tensor, delta=1.0)`

Reference: https://en.wikipedia.org/wiki/Huber_loss

Parameters

- **tensor** – (TensorFlow Tensor) the input value
- **delta** – (float) huber loss delta value

Returns (TensorFlow Tensor) huber loss output

`stable_baselines.common.tf_util.in_session (func)`

wrappes a function so that it is in a TensorFlow Session

Parameters **func** – (function) the function to wrap

Returns (function)

`stable_baselines.common.tf_util.initialize(sess=None)`

Initialize all the uninitialized variables in the global scope.

Parameters **sess** – (TensorFlow Session)

`stable_baselines.common.tf_util.intprod(tensor)`

calculates the product of all the elements in a list

Parameters **tensor** – ([Number]) the list of elements

Returns (int) the product truncated

`stable_baselines.common.tf_util.leaky_relu(tensor, leak=0.2)`

Leaky ReLU http://web.stanford.edu/~awni/papers/relu_hybrid_icml2013_final.pdf

Parameters

- **tensor** – (float) the input value
- **leak** – (float) the leaking coefficient when the function is saturated

Returns (float) Leaky ReLU output

`stable_baselines.common.tf_util.load_state(fname, sess=None, var_list=None)`

Load a TensorFlow saved model

Parameters

- **fname** – (str) the graph name
- **sess** – (TensorFlow Session) the session, if None: `get_default_session()`
- **var_list** – ([TensorFlow Tensor] or dict(str: TensorFlow Tensor)) A list of Variable/SaveableObject, or a dictionary mapping names to SaveableObject's. If None, defaults to the list of all saveable objects.

`stable_baselines.common.tf_util.make_session(num_cpu=None, make_default=False, graph=None)`

Returns a session that will use <num_cpu> CPU's only

Parameters

- **num_cpu** – (int) number of CPUs to use for TensorFlow
- **make_default** – (bool) if this should return an InteractiveSession or a normal Session
- **graph** – (TensorFlow Graph) the graph of the session

Returns (TensorFlow session)

`stable_baselines.common.tf_util.normc_initializer(std=1.0, axis=0)`

Return a parameter initializer for TensorFlow

Parameters

- **std** – (float) standard deviation
- **axis** – (int) the axis to normalize on

Returns (function)

`stable_baselines.common.tf_util.numel(tensor)`

get TensorFlow Tensor's number of elements

Parameters **tensor** – (TensorFlow Tensor) the input tensor

Returns (int) the number of elements

`stable_baselines.common.tf_util.outer_scope_getter(scope, new_scope="")`
 remove a scope layer for the getter

Parameters

- **scope** – (str) the layer to remove
- **new_scope** – (str) optional replacement name

Returns (function (function, str, *args, **kwargs): Tensorflow Tensor)

`stable_baselines.common.tf_util.save_state(fname, sess=None, var_list=None)`
 Save a TensorFlow model

Parameters

- **fname** – (str) the graph name
- **sess** – (TensorFlow Session) The tf session, if None, get_default_session()
- **var_list** – ([TensorFlow Tensor] or dict(str: TensorFlow Tensor)) A list of Variable/SaveableObject, or a dictionary mapping names to SaveableObject's. If None, defaults to the list of all saveable objects.

`stable_baselines.common.tf_util.single_threaded_session(make_default=False, graph=None)`

Returns a session which will only use a single CPU

Parameters

- **make_default** – (bool) if this should return an InteractiveSession or a normal Session
- **graph** – (TensorFlow Graph) the graph of the session

Returns (TensorFlow session)

`stable_baselines.common.tf_util.switch(condition, then_expression, else_expression)`
 Switches between two operations depending on a scalar value (int or bool). Note that both *then_expression* and *else_expression* should be symbolic tensors of the *same shape*.

Parameters

- **condition** – (TensorFlow Tensor) scalar tensor.
- **then_expression** – (TensorFlow Operation)
- **else_expression** – (TensorFlow Operation)

Returns (TensorFlow Operation) the switch output

`stable_baselines.common.tf_util.var_shape(tensor)`
 get TensorFlow Tensor shape

Parameters **tensor** – (TensorFlow Tensor) the input tensor

Returns ([int]) the shape

1.23 Command Utils

Helpers for scripts like `run_atari.py`.

`stable_baselines.common.cmd_util.arg_parser()`
 Create an empty `argparse.ArgumentParser`.

Returns (ArgumentParser)

`stable_baselines.common.cmd_util.atari_arg_parser()`
Create an `argparse.ArgumentParser` for `run_atari.py`.

Returns (ArgumentParser) parser {'-env': 'BreakoutNoFrameskip-v4', '-seed': 0, '-num-timesteps': int(1e7)}

`stable_baselines.common.cmd_util.make_atari_env(env_id, num_env, seed, wrapper_kwargs=None, start_index=0, allow_early_resets=True)`

Create a wrapped, monitored `SubprocVecEnv` for Atari.

Parameters

- **env_id** – (str) the environment ID
- **num_env** – (int) the number of environment you wish to have in subprocesses
- **seed** – (int) the initial seed for RNG
- **wrapper_kwargs** – (dict) the parameters for `wrap_deepmind` function
- **start_index** – (int) start rank index
- **allow_early_resets** – (bool) allows early reset of the environment

Returns (Gym Environment) The atari environment

`stable_baselines.common.cmd_util.make_mujoco_env(env_id, seed, allow_early_resets=True)`

Create a wrapped, monitored `gym.Env` for MuJoCo.

Parameters

- **env_id** – (str) the environment ID
- **seed** – (int) the initial seed for RNG
- **allow_early_resets** – (bool) allows early reset of the environment

Returns (Gym Environment) The mujoco environment

`stable_baselines.common.cmd_util.make_robotics_env(env_id, seed, rank=0, allow_early_resets=True)`

Create a wrapped, monitored `gym.Env` for MuJoCo.

Parameters

- **env_id** – (str) the environment ID
- **seed** – (int) the initial seed for RNG
- **rank** – (int) the rank of the environment (for logging)
- **allow_early_resets** – (bool) allows early reset of the environment

Returns (Gym Environment) The robotic environment

`stable_baselines.common.cmd_util.mujoco_arg_parser()`
Create an `argparse.ArgumentParser` for `run_mujoco.py`.

Returns (ArgumentParser) parser {'-env': 'Reacher-v2', '-seed': 0, '-num-timesteps': int(1e6), '-play': False}

`stable_baselines.common.cmd_util.robotics_arg_parser()`
Create an `argparse.ArgumentParser` for `run_mujoco.py`.

Returns (ArgumentParser) parser {‘-env’: ‘FetchReach-v0’, ‘-seed’: 0, ‘-num-timesteps’: int(1e6)}

1.24 Schedules

Schedules are used as hyperparameter for most of the algorithms, in order to change value of a parameter over time (usually the learning rate).

This file is used for specifying various schedules that evolve over time throughout the execution of the algorithm, such as:

- learning rate for the optimizer
- exploration epsilon for the epsilon greedy exploration strategy
- beta parameter for beta parameter in prioritized replay

Each schedule has a function *value(t)* which returns the current value of the parameter given the timestep *t* of the optimization procedure.

class `stable_baselines.common.schedules.ConstantSchedule` (*value*)

Value remains constant over time.

Parameters *value* – (float) Constant value of the schedule

value (*step*)

Value of the schedule for a given timestep

Parameters *step* – (int) the timestep

Returns (float) the output value for the given timestep

class `stable_baselines.common.schedules.LinearSchedule` (*schedule_timesteps*, *final_p*,
initial_p=1.0)

Linear interpolation between *initial_p* and *final_p* over *schedule_timesteps*. After this many timesteps pass *final_p* is returned.

Parameters

- ***schedule_timesteps*** – (int) Number of timesteps for which to linearly anneal *initial_p* to *final_p*
- ***initial_p*** – (float) initial output value
- ***final_p*** – (float) final output value

value (*step*)

Value of the schedule for a given timestep

Parameters *step* – (int) the timestep

Returns (float) the output value for the given timestep

class `stable_baselines.common.schedules.PiecewiseSchedule` (*endpoints*, *interpolation=<function linear_interpolation>*,
outside_value=None)

Piecewise schedule.

Parameters

- **endpoints** – ([int, int]) list of pairs (*time*, *value*) meaning that schedule should output *value* when $t == \text{time}$. All the values for time must be sorted in an increasing order. When t is between two times, e.g. (*time_a*, *value_a*) and (*time_b*, *value_b*), such that $\text{time}_a \leq t < \text{time}_b$ then value outputs *interpolation*(*value_a*, *value_b*, *alpha*) where alpha is a fraction of time passed between *time_a* and *time_b* for time t .
- **interpolation** – (lambda (float, float, float): float) a function that takes value to the left and to the right of t according to the *endpoints*. Alpha is the fraction of distance from left endpoint to right endpoint that t has covered. See *linear_interpolation* for example.
- **outside_value** – (float) if the value is requested outside of all the intervals specified in *endpoints* this value is returned. If None then AssertionError is raised when outside value is requested.

value (*step*)

Value of the schedule for a given timestep

Parameters *step* – (int) the timestep

Returns (float) the output value for the given timestep

`stable_baselines.common.schedules.linear_interpolation(left, right, alpha)`

Linear interpolation between *left* and *right*.

Parameters

- **left** – (float) left boundary
- **right** – (float) right boundary
- **alpha** – (float) coeff in $[0, 1]$

Returns (float)

1.25 Changelog

For download links, please look at [Github release page](#).

1.25.1 Release 2.2.0 (2018-11-07)

- Hotfix for ppo2, the wrong placeholder was used for the value function

1.25.2 Release 2.1.2 (2018-11-06)

- added `async_eigen_decomp` parameter for ACKTR and set it to `False` by default (remove deprecation warnings)
- added methods for calling env methods/setting attributes inside a `VecEnv` (thanks to @bjmuld)
- updated gym minimum version

1.25.3 Release 2.1.1 (2018-10-20)

- fixed MpiAdam synchronization issue in PPO1 (thanks to @brendenpetersen) issue #50
- fixed dependency issues (new mujoco-py requires a mujoco licence + gym broke MultiDiscrete space shape)

1.25.4 Release 2.1.0 (2018-10-2)

Warning: This version contains breaking changes for DQN policies, please read the full details

Bug fixes + doc update

- added patch fix for equal function using *gym.spaces.MultiDiscrete* and *gym.spaces.MultiBinary*
- fixes for DQN action_probability
- re-added double DQN + refactored DQN policies **breaking changes**
- replaced *async* with *async_eigen_decomp* in ACKTR/KFAC for python 3.7 compatibility
- removed action clipping for prediction of continuous actions (see issue #36)
- fixed NaN issue due to clipping the continuous action in the wrong place (issue #36)
- documentation was updated (policy + DDPG example hyperparameters)

1.25.5 Release 2.0.0 (2018-09-18)

Warning: This version contains breaking changes, please read the full details

Tensorboard, refactoring and bug fixes

- Renamed DeepQ to DQN **breaking changes**
- Renamed DeepQPolicy to DQNPolicy **breaking changes**
- fixed DDPG behavior **breaking changes**
- changed default policies for DDPG, so that DDPG now works correctly **breaking changes**
- added more documentation (some modules from common).
- added doc about using custom env
- added Tensorboard support for A2C, ACER, ACKTR, DDPG, DeepQ, PPO1, PPO2 and TRPO
- added episode reward to Tensorboard
- added documentation for Tensorboard usage
- added Identity for Box action space
- fixed render function ignoring parameters when using wrapped environments
- fixed PPO1 and TRPO done values for recurrent policies
- fixed image normalization not occurring when using images
- updated VecEnv objects for the new Gym version
- added test for DDPG
- refactored DQN policies
- added registry for policies, can be passed as string to the agent
- added documentation for custom policies + policy registration

- fixed numpy warning when using DDPG Memory
- fixed DummyVecEnv not copying the observation array when stepping and resetting
- added pre-built docker images + installation instructions
- added `deterministic` argument in the predict function
- added assert in PPO2 for recurrent policies
- fixed predict function to handle both vectorized and unwrapped environment
- added input check to the predict function
- refactored ActorCritic models to reduce code duplication
- refactored Off Policy models (to begin HER and replay_buffer refactoring)
- added tests for auto vectorization detection
- fixed render function, to handle positional arguments

1.25.6 Release 1.0.7 (2018-08-29)

Bug fixes and documentation

- added html documentation using sphinx + integration with read the docs
- cleaned up README + typos
- fixed normalization for DQN with images
- fixed DQN identity test

1.25.7 Release 1.0.1 (2018-08-20)

Refactored Stable Baselines

- refactored A2C, ACER, ACTKR, DDPG, DeepQ, GAIL, TRPO, PPO1 and PPO2 under a single constant class
- added callback to refactored algorithm training
- added saving and loading to refactored algorithms
- refactored ACER, DDPG, GAIL, PPO1 and TRPO to fit with A2C, PPO2 and ACKTR policies
- added new policies for most algorithms (Mlp, MlpLstm, MlpLnLstm, Cnn, CnnLstm and CnnLnLstm)
- added dynamic environment switching (so continual RL learning is now feasible)
- added prediction from observation and action probability from observation for all the algorithms
- fixed graphs issues, so models wont collide in names
- fixed behavior_clone weight loading for GAIL
- fixed Tensorflow using all the GPU VRAM
- fixed models so that they are all compatible with vectorized environments
- fixed ``set_global_seed`` to update ``gym.spaces``'s random seed
- fixed PPO1 and TRPO performance issues when learning identity function
- added new tests for loading, saving, continuous actions and learning the identity function
- fixed DQN wrapping for atari

- added saving and loading for Vecnormalize wrapper
- added automatic detection of action space (for the policy network)
- fixed ACER buffer with constant values assuming `n_stack=4`
- fixed some RL algorithms not clipping the action to be in the `action_space`, when using ``gym.spaces.Box``
- refactored algorithms can take either a ``gym.Environment`` or a ``str`` ([if the environment name is registered])(<https://github.com/openai/gym/wiki/Environments>))
- Hoftix in ACER (compared to v1.0.0)

Future Work :

- Finish refactoring HER
- Refactor ACKTR and ACER for continuous implementation

1.25.8 Release 0.1.6 (2018-07-27)

Deobfuscation of the code base + pep8 and fixes

- Fixed `tf.session().__enter__()` being used, rather than `sess = tf.session()` and passing the session to the objects
- Fixed uneven scoping of TensorFlow Sessions throughout the code
- Fixed rolling vecwrapper to handle observations that are not only grayscale images
- Fixed deepq saving the environment when trying to save itself
- Fixed `ValueError: Cannot take the length of Shape with unknown rank.` in `acktr`, when running `run_atari.py` script.
- Fixed calling baselines sequentially no longer creates graph conflicts
- Fixed mean on empty array warning with deepq
- Fixed `kfac` eigen decomposition not cast to `float64`, when the parameter `use_float64` is set to `True`
- Fixed Dataset data loader, not correctly resetting id position if shuffling is disabled
- Fixed `EOFError` when reading from connection in the `worker` in `subproc_vec_env.py`
- Fixed `behavior_clone` weight loading and saving for GAIL
- Avoid taking root square of negative number in `trpo_mpi.py`
- Removed some duplicated code (`a2cpolicy`, `trpo_mpi`)
- Removed unused, undocumented and crashing function `reset_task` in `subproc_vec_env.py`
- Reformatted code to PEP8 style
- Documented all the codebase
- Added atari tests
- Added logger tests

Missing: tests for `acktr` continuous (+ HER, gail but they rely on mujoco...)

1.25.9 Maintainers

Stable-Baselines is currently maintained by [Ashley Hill](#) (aka [@hill-a](#)) and [Antonin Raffin](#) (aka [@araffin](#)).

1.25.10 Contributors (since v2.0.0):

In random order. . .

Thanks to @bjmULD @iambenzo @iandanforth @r7vme @brendenpetersen @huvar

1.26 Plotting Results

`stable_baselines.results_plotter.main()`

Example usage in jupyter-notebook

```
from stable_baselines import log_viewer
%matplotlib inline
log_viewer.plot_results(["./log"], 10e6, log_viewer.X_TIMESTEPS, "Breakout")
```

Here `./log` is a directory containing the `monitor.csv` files

`stable_baselines.results_plotter.plot_curves(xy_list, xaxis, title)`

plot the curves

Parameters

- **xy_list** – ((np.ndarray, np.ndarray)) the x and y coordinates to plot
- **xaxis** – (str) the axis for the x and y output (can be `X_TIMESTEPS='timesteps'`, `X_EPISODES='episodes'` or `X_WALLTIME='walltime_hrs'`)
- **title** – (str) the title of the plot

`stable_baselines.results_plotter.plot_results(dirs, num_timesteps, xaxis, task_name)`

plot the results

Parameters

- **dirs** – ([str]) the save location of the results to plot
- **num_timesteps** – (int) only plot the points below this value
- **xaxis** – (str) the axis for the x and y output (can be `X_TIMESTEPS='timesteps'`, `X_EPISODES='episodes'` or `X_WALLTIME='walltime_hrs'`)
- **task_name** – (str) the title of the task to plot

`stable_baselines.results_plotter.rolling_window(array, window)`

apply a rolling window to a np.ndarray

Parameters

- **array** – (np.ndarray) the input Array
- **window** – (int) length of the rolling window

Returns (np.ndarray) rolling window on the input array

`stable_baselines.results_plotter.ts2xy(timesteps, xaxis)`

Decompose a timesteps variable to x and y

Parameters

- **timesteps** – (Pandas DataFrame) the input data
- **xaxis** – (str) the axis for the x and y output (can be `X_TIMESTEPS='timesteps'`, `X_EPISODES='episodes'` or `X_WALLTIME='walltime_hrs'`)

Returns (np.ndarray, np.ndarray) the x and y output

`stable_baselines.results_plotter.window_func(var_1, var_2, window, func)`
apply a function to the rolling window of 2 arrays

Parameters

- **var_1** – (np.ndarray) variable 1
- **var_2** – (np.ndarray) variable 2
- **window** – (int) length of the rolling window
- **func** – (numpy function) function to apply on the rolling window on variable 2 (such as `np.mean`)

Returns (np.ndarray, np.ndarray) the rolling output with applied function

Citing Stable Baselines

To cite this project in publications:

```
@misc{stable-baselines,  
  author = {Hill, Ashley and Raffin, Antonin and Traore, Rene and Dhariwal, Prafulla,  
↪and Hesse, Christopher and Klimov, Oleg and Nichol, Alex and Plappert, Matthias and,  
↪Radford, Alec and Schulman, John and Sidor, Szymon and Wu, Yuhuai},  
  title = {Stable Baselines},  
  year = {2018},  
  publisher = {GitHub},  
  journal = {GitHub repository},  
  howpublished = {\url{https://github.com/hill-a/stable-baselines}},  
}
```


CHAPTER 3

Contributing

To any interested in making the rl baselines better, there is still some improvements that needs to be done: good-to-have features like support for continuous actions (ACER) and more documentation on the rl algorithms.

If you want to contribute, please open an issue first and then propose your pull request on Github at <https://github.com/hill-a/stable-baselines>.

CHAPTER 4

Indices and tables

- `genindex`
- `search`
- `modindex`

S

`stable_baselines.a2c`, 30
`stable_baselines.acer`, 33
`stable_baselines.acktr`, 36
`stable_baselines.common.base_class`, 22
`stable_baselines.common.cmd_util`, 82
`stable_baselines.common.distributions`,
72
`stable_baselines.common.policies`, 24
`stable_baselines.common.schedules`, 84
`stable_baselines.common.tf_util`, 79
`stable_baselines.common.vec_env`, 13
`stable_baselines.ddpg`, 39
`stable_baselines.deepq`, 50
`stable_baselines.gail`, 57
`stable_baselines.her`, 61
`stable_baselines.ppo1`, 62
`stable_baselines.ppo2`, 66
`stable_baselines.results_plotter`, 89
`stable_baselines.trpo_mpi`, 69

A

A2C (class in `stable_baselines.a2c`), 31
 ACER (class in `stable_baselines.acer`), 34
 ACKTR (class in `stable_baselines.acktr`), 37
`action_probability()` (`stable_baselines.a2c.A2C` method), 32
`action_probability()` (`stable_baselines.acer.ACER` method), 35
`action_probability()` (`stable_baselines.acktr.ACKTR` method), 38
`action_probability()` (`stable_baselines.common.base_class.BaseRLModel` method), 23
`action_probability()` (`stable_baselines.ddpg.DDPG` method), 42
`action_probability()` (`stable_baselines.deepq.DQN` method), 53
`action_probability()` (`stable_baselines.gail.GAIL` method), 59
`action_probability()` (`stable_baselines.her.HER` method), 61
`action_probability()` (`stable_baselines.ppo1.PPO1` method), 65
`action_probability()` (`stable_baselines.ppo2.PPO2` method), 68
`action_probability()` (`stable_baselines.trpo_mpi.TRPO` method), 70
`ActorCriticPolicy` (class in `stable_baselines.common.policies`), 25
`adapt()` (`stable_baselines.ddpg.AdaptiveParamNoiseSpec` method), 49
`AdaptiveParamNoiseSpec` (class in `stable_baselines.ddpg`), 48
`arg_parser()` (in module `stable_baselines.common.cmd_util`), 82
`atari_arg_parser()` (in module `stable_baselines.common.cmd_util`), 83

B

`BaseRLModel` (class in `stable_baselines.common.base_class`), 22
`BernoulliProbabilityDistribution` (class in `stable_baselines.common.distributions`), 72
`BernoulliProbabilityDistributionType` (class in `stable_baselines.common.distributions`), 72

C

`CategoricalProbabilityDistribution` (class in `stable_baselines.common.distributions`), 73
`CategoricalProbabilityDistributionType` (class in `stable_baselines.common.distributions`), 74
`close()` (`stable_baselines.common.vec_env.DummyVecEnv` method), 14
`close()` (`stable_baselines.common.vec_env.SubprocVecEnv` method), 15
`close()` (`stable_baselines.common.vec_env.VecFrameStack` method), 16
`CnnLnLstmPolicy` (class in `stable_baselines.common.policies`), 30
`CnnLstmPolicy` (class in `stable_baselines.common.policies`), 29
`CnnPolicy` (class in `stable_baselines.common.policies`), 29
`CnnPolicy` (class in `stable_baselines.ddpg`), 46
`CnnPolicy` (class in `stable_baselines.deepq`), 55
`ConstantSchedule` (class in `stable_baselines.common.schedules`), 84
`conv2d()` (in module `stable_baselines.common.tf_util`), 79

D

`DDPG` (class in `stable_baselines.ddpg`), 41
`DiagGaussianProbabilityDistribution` (class in `stable_baselines.common.distributions`), 74
`DiagGaussianProbabilityDistributionType` (class in `stable_baselines.common.distributions`), 75
`display_var_info()` (in module `stable_baselines.common.tf_util`), 79

DQN (class in `stable_baselines.deepq`), [52](#)
 DummyVecEnv (class in `stable_baselines.common.vec_env`), [14](#)

E

`entropy()` (`stable_baselines.common.distributions.BernoulliProbabilityDistribution` method), [72](#)
`entropy()` (`stable_baselines.common.distributions.CategoricalProbabilityDistribution` method), [73](#)
`entropy()` (`stable_baselines.common.distributions.DiagGaussianProbabilityDistribution` method), [74](#)
`entropy()` (`stable_baselines.common.distributions.MultiCategoricalProbabilityDistribution` method), [76](#)
`entropy()` (`stable_baselines.common.distributions.ProbabilityDistribution` method), [77](#)
`env_method()` (`stable_baselines.common.vec_env.DummyVecEnv` method), [14](#)
`env_method()` (`stable_baselines.common.vec_env.SubprocVecEnv` method), [15](#)

F

`FeedForwardPolicy` (class in `stable_baselines.common.policies`), [26](#)
`flatgrad()` (in module `stable_baselines.common.tf_util`), [79](#)
`flatparam()` (`stable_baselines.common.distributions.BernoulliProbabilityDistribution` method), [72](#)
`flatparam()` (`stable_baselines.common.distributions.CategoricalProbabilityDistribution` method), [73](#)
`flatparam()` (`stable_baselines.common.distributions.DiagGaussianProbabilityDistribution` method), [74](#)
`flatparam()` (`stable_baselines.common.distributions.MultiCategoricalProbabilityDistribution` method), [76](#)
`flatparam()` (`stable_baselines.common.distributions.ProbabilityDistribution` method), [77](#)
`flattenallbut0()` (in module `stable_baselines.common.tf_util`), [79](#)

`fromflat()` (`stable_baselines.common.distributions.BernoulliProbabilityDistribution` class method), [72](#)
`fromflat()` (`stable_baselines.common.distributions.CategoricalProbabilityDistribution` class method), [73](#)
`fromflat()` (`stable_baselines.common.distributions.DiagGaussianProbabilityDistribution` class method), [74](#)
`fromflat()` (`stable_baselines.common.distributions.MultiCategoricalProbabilityDistribution` class method), [76](#)
`function()` (in module `stable_baselines.common.tf_util`), [80](#)

G

GAIL (class in `stable_baselines.gail`), [59](#)
`get_attr()` (`stable_baselines.common.vec_env.DummyVecEnv` method), [14](#)
`get_attr()` (`stable_baselines.common.vec_env.SubprocVecEnv` method), [15](#)

`get_available_gpus()` (in module `stable_baselines.common.tf_util`), [80](#)
`get_env()` (`stable_baselines.a2c.A2C` method), [32](#)
`get_env()` (`stable_baselines.acer.ACER` method), [35](#)
`get_env()` (`stable_baselines.acktr.ACKTR` method), [38](#)
`get_env()` (`stable_baselines.common.base_class.BaseRLModel` method), [23](#)
`get_env()` (`stable_baselines.ddpg.DDPG` method), [42](#)
`get_env()` (`stable_baselines.deepq.DQN` method), [53](#)
`get_env()` (`stable_baselines.gail.GAIL` method), [60](#)
`get_env()` (`stable_baselines.ppo1.PPO1` method), [65](#)
`get_env()` (`stable_baselines.ppo2.PPO2` method), [68](#)
`get_env()` (`stable_baselines.trpo_mpi.TRPO` method), [70](#)
`get_trainable_vars()` (in module `stable_baselines.common.tf_util`), [80](#)
`get_images()` (`stable_baselines.common.vec_env.DummyVecEnv` method), [14](#)
`get_images()` (`stable_baselines.common.vec_env.SubprocVecEnv` method), [15](#)
`get_original_obs()` (`stable_baselines.common.vec_env.VecNormalize` method), [17](#)
`get_stats()` (`stable_baselines.ddpg.AdaptiveParamNoiseSpec` method), [49](#)
`get_trainable_vars()` (in module `stable_baselines.common.tf_util`), [80](#)

H

HER (class in `stable_baselines.her`), [61](#)
`huber_loss()` (in module `stable_baselines.common.tf_util`), [80](#)
`init()` (in module `stable_baselines.common.tf_util`), [80](#)
`initialize()` (in module `stable_baselines.common.tf_util`), [81](#)
`insert0()` (in module `stable_baselines.common.tf_util`), [81](#)

K

`kl()` (`stable_baselines.common.distributions.BernoulliProbabilityDistribution` method), [72](#)
`kl()` (`stable_baselines.common.distributions.CategoricalProbabilityDistribution` method), [73](#)
`kl()` (`stable_baselines.common.distributions.DiagGaussianProbabilityDistribution` method), [74](#)
`kl()` (`stable_baselines.common.distributions.MultiCategoricalProbabilityDistribution` method), [76](#)
`kl()` (`stable_baselines.common.distributions.ProbabilityDistribution` method), [77](#)

L

`leaky_relu()` (in module `stable_baselines.common.tf_util`), [81](#)

[learn\(\)](#) (stable_baselines.a2c.A2C method), [32](#)
[learn\(\)](#) (stable_baselines.acer.ACER method), [35](#)
[learn\(\)](#) (stable_baselines.acktr.ACKTR method), [38](#)
[learn\(\)](#) (stable_baselines.common.base_class.BaseRLModel method), [23](#)
[learn\(\)](#) (stable_baselines.ddpg.DDPG method), [42](#)
[learn\(\)](#) (stable_baselines.deepq.DQN method), [53](#)
[learn\(\)](#) (stable_baselines.gail.GAIL method), [60](#)
[learn\(\)](#) (stable_baselines.her.HER method), [62](#)
[learn\(\)](#) (stable_baselines.ppo1.PPO1 method), [65](#)
[learn\(\)](#) (stable_baselines.ppo2.PPO2 method), [68](#)
[learn\(\)](#) (stable_baselines.trpo_mpi.TRPO method), [70](#)
[linear_interpolation\(\)](#) (in module `stable_baselines.common.schedules`), [85](#)
[LinearSchedule](#) (class in `stable_baselines.common.schedules`), [84](#)
[LnCnnPolicy](#) (class in `stable_baselines.ddpg`), [47](#)
[LnCnnPolicy](#) (class in `stable_baselines.deepq`), [56](#)
[LnMlpPolicy](#) (class in `stable_baselines.ddpg`), [44](#)
[LnMlpPolicy](#) (class in `stable_baselines.deepq`), [55](#)
[load\(\)](#) (stable_baselines.a2c.A2C class method), [32](#)
[load\(\)](#) (stable_baselines.acer.ACER class method), [35](#)
[load\(\)](#) (stable_baselines.acktr.ACKTR class method), [38](#)
[load\(\)](#) (stable_baselines.common.base_class.BaseRLModel class method), [23](#)
[load\(\)](#) (stable_baselines.ddpg.DDPG class method), [42](#)
[load\(\)](#) (stable_baselines.deepq.DQN class method), [53](#)
[load\(\)](#) (stable_baselines.gail.GAIL class method), [60](#)
[load\(\)](#) (stable_baselines.her.HER class method), [62](#)
[load\(\)](#) (stable_baselines.ppo1.PPO1 class method), [65](#)
[load\(\)](#) (stable_baselines.ppo2.PPO2 class method), [68](#)
[load\(\)](#) (stable_baselines.trpo_mpi.TRPO class method), [71](#)
[load_running_average\(\)](#) (stable_baselines.common.vec_env.VecNormalize method), [17](#)
[load_state\(\)](#) (in module `stable_baselines.common.tf_util`), [81](#)
[logp\(\)](#) (stable_baselines.common.distributions.ProbabilityDistribution method), [77](#)
[LstmPolicy](#) (class in `stable_baselines.common.policies`), [27](#)

M

[main\(\)](#) (in module `stable_baselines.results_plotter`), [89](#)
[make_actor\(\)](#) (stable_baselines.ddpg.CnnPolicy method), [46](#)
[make_actor\(\)](#) (stable_baselines.ddpg.LnCnnPolicy method), [47](#)
[make_actor\(\)](#) (stable_baselines.ddpg.LnMlpPolicy method), [45](#)
[make_actor\(\)](#) (stable_baselines.ddpg.MlpPolicy method), [43](#)

[make_atari_env\(\)](#) (in module `stable_baselines.common.cmd_util`), [83](#)
[make_critic\(\)](#) (stable_baselines.ddpg.CnnPolicy method), [46](#)
[make_critic\(\)](#) (stable_baselines.ddpg.LnCnnPolicy method), [47](#)
[make_critic\(\)](#) (stable_baselines.ddpg.LnMlpPolicy method), [45](#)
[make_critic\(\)](#) (stable_baselines.ddpg.MlpPolicy method), [43](#)
[make_mujoco_env\(\)](#) (in module `stable_baselines.common.cmd_util`), [83](#)
[make_proba_dist_type\(\)](#) (in module `stable_baselines.common.distributions`), [78](#)
[make_robotics_env\(\)](#) (in module `stable_baselines.common.cmd_util`), [83](#)
[make_session\(\)](#) (in module `stable_baselines.common.tf_util`), [81](#)
[MlpLnLstmPolicy](#) (class in `stable_baselines.common.policies`), [28](#)
[MlpLstmPolicy](#) (class in `stable_baselines.common.policies`), [28](#)
[MlpPolicy](#) (class in `stable_baselines.common.policies`), [28](#)
[MlpPolicy](#) (class in `stable_baselines.ddpg`), [43](#)
[MlpPolicy](#) (class in `stable_baselines.deepq`), [54](#)
[mode\(\)](#) (stable_baselines.common.distributions.BernoulliProbabilityDistribution method), [72](#)
[mode\(\)](#) (stable_baselines.common.distributions.CategoricalProbabilityDistribution method), [73](#)
[mode\(\)](#) (stable_baselines.common.distributions.DiagGaussianProbabilityDistribution method), [75](#)
[mode\(\)](#) (stable_baselines.common.distributions.MultiCategoricalProbabilityDistribution method), [76](#)
[mode\(\)](#) (stable_baselines.common.distributions.ProbabilityDistribution method), [77](#)
[mujoco_arg_parser\(\)](#) (in module `stable_baselines.common.cmd_util`), [83](#)
[MultiCategoricalProbabilityDistribution](#) (class in `stable_baselines.common.distributions`), [75](#)
[MultiCategoricalProbabilityDistributionType](#) (class in `stable_baselines.common.distributions`), [76](#)

N

[neglogp\(\)](#) (stable_baselines.common.distributions.BernoulliProbabilityDistribution method), [72](#)
[neglogp\(\)](#) (stable_baselines.common.distributions.CategoricalProbabilityDistribution method), [73](#)
[neglogp\(\)](#) (stable_baselines.common.distributions.DiagGaussianProbabilityDistribution method), [75](#)
[neglogp\(\)](#) (stable_baselines.common.distributions.MultiCategoricalProbabilityDistribution method), [76](#)
[neglogp\(\)](#) (stable_baselines.common.distributions.ProbabilityDistribution method), [77](#)

NormalActionNoise (class in stable_baselines.ddpg), 49
 normc_initializer() (in module stable_baselines.common.tf_util), 81
 numel() (in module stable_baselines.common.tf_util), 81

O

OrnsteinUhlenbeckActionNoise (class in stable_baselines.ddpg), 49
 outer_scope_getter() (in module stable_baselines.common.tf_util), 82

P

param_placeholder() (stable_baselines.common.distributions.ProbabilityDistributionType method), 78
 param_shape() (stable_baselines.common.distributions.BernoulliProbabilityDistributionType method), 72
 param_shape() (stable_baselines.common.distributions.CategoricalProbabilityDistributionType method), 74
 param_shape() (stable_baselines.common.distributions.DiagGaussianProbabilityDistributionType method), 75
 param_shape() (stable_baselines.common.distributions.MultiCategoricalProbabilityDistributionType method), 76
 param_shape() (stable_baselines.common.distributions.ProbabilityDistributionType method), 78
 PiecewiseSchedule (class in stable_baselines.common.schedules), 84
 plot_curves() (in module stable_baselines.results_plotter), 89
 plot_results() (in module stable_baselines.results_plotter), 89
 PPO1 (class in stable_baselines.ppo1), 64
 PPO2 (class in stable_baselines.ppo2), 67
 predict() (stable_baselines.a2c.A2C method), 32
 predict() (stable_baselines.acer.ACER method), 36
 predict() (stable_baselines.acktr.ACKTR method), 39
 predict() (stable_baselines.common.base_class.BaseRLModel method), 23
 predict() (stable_baselines.ddpg.DDPG method), 42
 predict() (stable_baselines.deepq.DQN method), 53
 predict() (stable_baselines.gail.GAIL method), 60
 predict() (stable_baselines.her.HER method), 62
 predict() (stable_baselines.ppo1.PPO1 method), 65
 predict() (stable_baselines.ppo2.PPO2 method), 69
 predict() (stable_baselines.trpo_mpi.TRPO method), 71
 proba_distribution_from_flat() (stable_baselines.common.distributions.DiagGaussianProbabilityDistributionType method), 75
 proba_distribution_from_flat() (stable_baselines.common.distributions.MultiCategoricalProbabilityDistributionType method), 76
 proba_distribution_from_flat() (stable_baselines.common.distributions.ProbabilityDistributionType method), 78

proba_distribution_from_latent() (stable_baselines.common.distributions.BernoulliProbabilityDistributionType method), 72
 proba_distribution_from_latent() (stable_baselines.common.distributions.CategoricalProbabilityDistributionType method), 74
 proba_distribution_from_latent() (stable_baselines.common.distributions.DiagGaussianProbabilityDistributionType method), 75
 proba_distribution_from_latent() (stable_baselines.common.distributions.MultiCategoricalProbabilityDistributionType method), 76
 proba_distribution_from_latent() (stable_baselines.common.distributions.ProbabilityDistributionType method), 78
 proba_step() (stable_baselines.common.policies.ActorCriticPolicy method), 25
 proba_step() (stable_baselines.common.policies.FeedForwardPolicy method), 26
 proba_step() (stable_baselines.common.policies.LstmPolicy method), 27
 proba_step() (stable_baselines.ddpg.CnnPolicy method), 46
 proba_step() (stable_baselines.ddpg.LnCnnPolicy method), 48
 proba_step() (stable_baselines.ddpg.LnMlpPolicy method), 45
 proba_step() (stable_baselines.ddpg.MlpPolicy method), 44
 proba_step() (stable_baselines.deepq.CnnPolicy method), 56
 proba_step() (stable_baselines.deepq.LnCnnPolicy method), 57
 proba_step() (stable_baselines.deepq.LnMlpPolicy method), 55
 proba_step() (stable_baselines.deepq.MlpPolicy method), 54
 probability_distribution_class() (stable_baselines.common.distributions.BernoulliProbabilityDistributionType method), 73
 probability_distribution_class() (stable_baselines.common.distributions.CategoricalProbabilityDistributionType method), 74
 probability_distribution_class() (stable_baselines.common.distributions.DiagGaussianProbabilityDistributionType method), 75
 probability_distribution_class() (stable_baselines.common.distributions.MultiCategoricalProbabilityDistributionType method), 77
 probability_distribution_class() (stable_baselines.common.distributions.ProbabilityDistributionType method), 78
 ProbabilityDistribution (class in stable_baselines.common.distributions), 77

ProbabilityDistributionType (class in stable_baselines.common.distributions), 77

R

render() (stable_baselines.common.vec_env.DummyVecEnv method), 14

render() (stable_baselines.common.vec_env.SubprocVecEnv method), 15

reset() (stable_baselines.common.vec_env.DummyVecEnv method), 14

reset() (stable_baselines.common.vec_env.SubprocVecEnv method), 15

reset() (stable_baselines.common.vec_env.VecFrameStack method), 16

reset() (stable_baselines.common.vec_env.VecNormalize method), 17

reset() (stable_baselines.ddpg.NormalActionNoise method), 49

reset() (stable_baselines.ddpg.OrnsteinUhlenbeckActionNoise method), 49

robotics_arg_parser() (in module stable_baselines.common.cmd_util), 83

rolling_window() (in module stable_baselines.results_plotter), 89

S

sample() (stable_baselines.common.distributions.BernoulliProbabilityDistribution method), 72

sample() (stable_baselines.common.distributions.CategoricalProbabilityDistribution method), 74

sample() (stable_baselines.common.distributions.DiagGaussianProbabilityDistribution method), 75

sample() (stable_baselines.common.distributions.MultiCategoricalProbabilityDistribution method), 76

sample() (stable_baselines.common.distributions.ProbabilityDistribution method), 77

sample_dtype() (stable_baselines.common.distributions.BernoulliProbabilityDistributionType method), 73

sample_dtype() (stable_baselines.common.distributions.CategoricalProbabilityDistributionType method), 74

sample_dtype() (stable_baselines.common.distributions.DiagGaussianProbabilityDistributionType method), 75

sample_dtype() (stable_baselines.common.distributions.MultiCategoricalProbabilityDistributionType method), 77

sample_dtype() (stable_baselines.common.distributions.ProbabilityDistributionType method), 78

sample_placeholder() (stable_baselines.common.distributions.ProbabilityDistributionType method), 78

sample_shape() (stable_baselines.common.distributions.BernoulliProbabilityDistributionType method), 73

sample_shape() (stable_baselines.common.distributions.CategoricalProbabilityDistributionType method), 74

sample_shape() (stable_baselines.common.distributions.DiagGaussianProbabilityDistributionType method), 75

sample_shape() (stable_baselines.common.distributions.MultiCategoricalProbabilityDistributionType method), 77

sample_shape() (stable_baselines.common.distributions.ProbabilityDistributionType method), 78

save() (stable_baselines.a2c.A2C method), 33

save() (stable_baselines.acer.ACER method), 36

save() (stable_baselines.acktr.ACKTR method), 39

save() (stable_baselines.common.base_class.BaseRLModel method), 24

save() (stable_baselines.ddpg.DDPG method), 43

save() (stable_baselines.deepq.DQN method), 54

save() (stable_baselines.gail.GAIL method), 60

save() (stable_baselines.her.HER method), 62

save() (stable_baselines.ppo1.PPO1 method), 65

save() (stable_baselines.ppo2.PPO2 method), 69

save() (stable_baselines.trpo_mpi.TRPO method), 71

save_running_average() (stable_baselines.common.vec_env.VecNormalize method), 17

save_state() (in module stable_baselines.common.tf_util), 82

set_attr() (stable_baselines.common.vec_env.DummyVecEnv method), 14

set_attr() (stable_baselines.common.vec_env.SubprocVecEnv method), 15

set_env() (stable_baselines.a2c.A2C method), 33

set_env() (stable_baselines.acer.ACER method), 36

set_env() (stable_baselines.acktr.ACKTR method), 39

set_env() (stable_baselines.common.base_class.BaseRLModel method), 24

set_env() (stable_baselines.ddpg.DDPG method), 43

set_env() (stable_baselines.deepq.DQN method), 54

set_env() (stable_baselines.gail.GAIL method), 60

set_env() (stable_baselines.ppo1.PPO1 method), 66

set_env() (stable_baselines.ppo2.PPO2 method), 69

set_env() (stable_baselines.trpo_mpi.TRPO method), 71

setup_model() (stable_baselines.a2c.A2C method), 33

setup_model() (stable_baselines.acer.ACER method), 36

setup_model() (stable_baselines.acktr.ACKTR method), 39

setup_model() (stable_baselines.common.base_class.BaseRLModel method), 24

setup_model() (stable_baselines.ddpg.DDPG method), 43

setup_model() (stable_baselines.deepq.DQN method), 54

setup_model() (stable_baselines.gail.GAIL method), 61

setup_model() (stable_baselines.her.HER method), 62

setup_model() (stable_baselines.ppo1.PPO1 method), 66

setup_model() (stable_baselines.ppo2.PPO2 method), 69

setup_model() (stable_baselines.trpo_mpi.TRPO method), 71

`shape_el()` (in module `stable_baselines.common.distributions`), 79
`single_threaded_session()` (in module `stable_baselines.common.tf_util`), 82
`stable_baselines.a2c` (module), 30
`stable_baselines.acer` (module), 33
`stable_baselines.acktr` (module), 36
`stable_baselines.common.base_class` (module), 22
`stable_baselines.common.cmd_util` (module), 82
`stable_baselines.common.distributions` (module), 72
`stable_baselines.common.policies` (module), 24
`stable_baselines.common.schedules` (module), 84
`stable_baselines.common.tf_util` (module), 79
`stable_baselines.common.vec_env` (module), 13
`stable_baselines.ddpg` (module), 39
`stable_baselines.deepq` (module), 50
`stable_baselines.gail` (module), 57
`stable_baselines.her` (module), 61
`stable_baselines.ppo1` (module), 62
`stable_baselines.ppo2` (module), 66
`stable_baselines.results_plotter` (module), 89
`stable_baselines.trpo_mpi` (module), 69
`step()` (`stable_baselines.common.policies.ActorCriticPolicy` method), 25
`step()` (`stable_baselines.common.policies.FeedForwardPolicy` method), 26
`step()` (`stable_baselines.common.policies.LstmPolicy` method), 27
`step()` (`stable_baselines.ddpg.CnnPolicy` method), 47
`step()` (`stable_baselines.ddpg.LnCnnPolicy` method), 48
`step()` (`stable_baselines.ddpg.LnMlpPolicy` method), 45
`step()` (`stable_baselines.ddpg.MlpPolicy` method), 44
`step()` (`stable_baselines.deepq.CnnPolicy` method), 56
`step()` (`stable_baselines.deepq.LnCnnPolicy` method), 57
`step()` (`stable_baselines.deepq.LnMlpPolicy` method), 55
`step()` (`stable_baselines.deepq.MlpPolicy` method), 54
`step_async()` (`stable_baselines.common.vec_env.DummyVecEnv` method), 14
`step_async()` (`stable_baselines.common.vec_env.SubprocVecEnv` method), 16
`step_wait()` (`stable_baselines.common.vec_env.DummyVecEnv` method), 15
`step_wait()` (`stable_baselines.common.vec_env.SubprocVecEnv` method), 16
`step_wait()` (`stable_baselines.common.vec_env.VecFrameStack` method), 16
`step_wait()` (`stable_baselines.common.vec_env.VecNormalize` method), 17
`SubprocVecEnv` (class in `stable_baselines.common.vec_env`), 15
`switch()` (in module `stable_baselines.common.tf_util`), 82

T

TRPO (class in `stable_baselines.trpo_mpi`), 70

`ts2xy()` (in module `stable_baselines.results_plotter`), 89

V

`value()` (`stable_baselines.common.policies.ActorCriticPolicy` method), 25
`value()` (`stable_baselines.common.policies.FeedForwardPolicy` method), 26
`value()` (`stable_baselines.common.policies.LstmPolicy` method), 28
`value()` (`stable_baselines.common.schedules.ConstantSchedule` method), 84
`value()` (`stable_baselines.common.schedules.LinearSchedule` method), 84
`value()` (`stable_baselines.common.schedules.PiecewiseSchedule` method), 85
`value()` (`stable_baselines.ddpg.CnnPolicy` method), 47
`value()` (`stable_baselines.ddpg.LnCnnPolicy` method), 48
`value()` (`stable_baselines.ddpg.LnMlpPolicy` method), 45
`value()` (`stable_baselines.ddpg.MlpPolicy` method), 44
`var_shape()` (in module `stable_baselines.common.tf_util`), 82
`VecFrameStack` (class in `stable_baselines.common.vec_env`), 16
`VecNormalize` (class in `stable_baselines.common.vec_env`), 16

W

`window_func()` (in module `stable_baselines.results_plotter`), 90