

Setup

Setup instructions

This training depends on `kubectl`, the Kubernetes command-line interface.

Follow the instructions on the subsequent pages to complete the setup on your platform of choice.

Warning

In case you've already installed `kubectl`, please make sure you have an up-to-date version.

1. Local usage

Please follow the instructions on the *1.1. cli installation* page to install `kubect1` .

1.1. cli installation

The `kubect1` command is the command-line interface to work with one or several Kubernetes clusters.

The client is written in Go and you can run the single binary on the following operating systems:

- *1.1.1. Windows*
- *1.1.2. macOS*
- *1.1.3. Linux*

1.2. Verification

Verify the installation

You should now be able to execute `kubectl` in the command prompt. To test, execute:

```
kubectl version
```

You should now see something like (the version number may vary):

```
Client Version: version.Info{Major:"1", Minor:"18", GitVersion:"v1.18.0", GitCommit:"9e991415386e4cf155a24b1da15becaa390438d8", GitTreeState:"clean", BuildDate:"2020-03-25T14:58:59Z", GoVersion:"go1.13.8", Compiler:"gc", Platform:"linux/amd64"}
...
```

If you don't see a similar output, possibly there are issues with the `PATH` variable.

Warning

Make sure to use at least version 1.16.x for your `kubectl`

First steps with kubectl

The `kubectl` binary has many subcommands. Invoke `kubectl --help` (or simply `-h`) to get a list of all subcommands; `kubectl <subcommand> --help` gives you detailed help about a subcommand.

Optional tools

Have a look at the optional tools described in *2. Optional Kubernetes power tools* if you're interested.

Next steps

When you're ready to go, head on over to the [labs](#) and begin with the training!

2. Optional Kubernetes power tools

Optional Kubernetes power tools

`kubectx` and `kubens` are two handy shell scripts which let you easily switch between Kubernetes contexts and namespaces. See <https://github.com/ahmetb/kubectx> for detailed instructions.

Installation of `kubectx` and `kubens` :

```
curl https://raw.githubusercontent.com/ahmetb/kubectx/master/kubectx -o ~/bin/kubectx
curl https://raw.githubusercontent.com/ahmetb/kubectx/master/kubens -o ~/bin/kubens
chmod +x ~/bin/kubectx ~/bin/kubens
```

`kube-ps1` is another helpful shell script which adds the current context and namespace to the shell prompt: <https://github.com/jonmosco/kube-ps1>

`fzf` is yet another handy helper tool when you have to deal with a lot of contexts or namespaces by adding an interactive menu to `kubectx` and `kubens` : <https://github.com/junegunn/fzf>

`stern` is a very powerful enhancement of `kubectl logs` and lets you tail logs of multiple containers and Pods at the same time: <https://github.com/wercker/stern> .

Other tools to work with Kubernetes

- <https://github.com/lensapp/lens>

Next steps

When you're ready to go, head on over to the [labs](#) and begin with the training!

Labs

The purpose of these labs is to convey Kubernetes basics by providing hands-on tasks for people. Kubernetes (K8s) will allow you to deploy and deliver your software packaged as containers in an easy, straightforward way.

Goals of these labs:

- Help you get started with this modern technology
- Explain the basic concepts to you
- Show you how to deploy your first applications on Kubernetes

Additional Docs

- [Kubernetes Docs](#)
- [Helm Docs](#)

Additional Tutorials

- [Official Kubernetes Tutorial](#)

1. Introduction

In this lab, we will introduce the core concepts of Kubernetes.

All explanations and resources used in this lab give only a quick and not detailed overview. Please check [the official documentation](#) to get further details.

Core concepts

With the open source software Kubernetes, you get a platform to deploy your application in a container and operate it at the same time. Therefore, Kubernetes is also called a *Container Platform*, or the term *Container-as-a-Service* (CaaS) is used.

Depending on the configuration the term *Platform-as-a-Service* (PaaS) works as well.

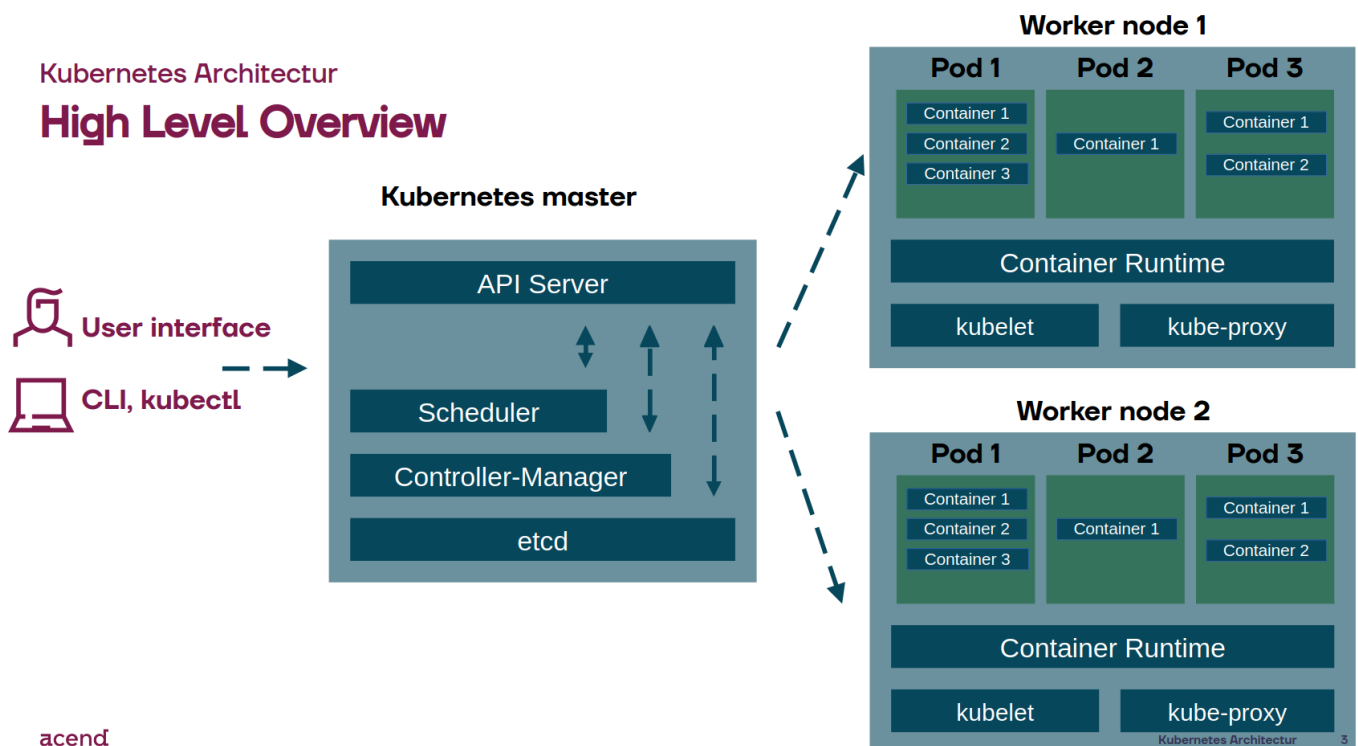
Container engine

Kubernetes' underlying container engine most often is [Docker](#). There are other container engines that could be used with Kubernetes such as [CRI-O](#).

Docker was originally created to help developers test their applications in their continuous integration environments. Nowadays, system admins also use it. CRI-O doesn't exist as long as Docker does. It is a "lightweight container runtime for Kubernetes" and is fully [OCI-compliant](#).

Overview

Kubernetes consists of control plane and worker (minion, compute) nodes.



Control plane and worker nodes

The control plane components are the *API server*, the *scheduler* and the *controller manager*. The API server itself represents the management interface. The scheduler and the controller manager decide how applications should be deployed on the cluster. Additionally, the state and configuration of the cluster itself are controlled in the control plane components.

Worker nodes are also known as compute nodes, application nodes or minions, and are responsible for running the container workload (applications). The *control plane* for the worker nodes is implemented in the control plane components. The hosts running these components were historically called masters.

Containers and images

The smallest entities in Kubernetes are Pods, which resemble your containerized application.

Using container virtualization, processes on a Linux system can be isolated up to a level where only the predefined resources are available. Several containers can run on the same system without “seeing” each other (files, process IDs, network). One container should contain one application (web server, database, cache, etc.). It should be at least one part of the application, e.g. when running a multi-service middleware. In a container itself any process can be started that runs natively on your operating system.

Containers are based on images. An image represents the file tree, which includes the binary, shared libraries and other files which are needed to run your application.

A container image is typically built from a `Containerfile` or `Dockerfile`, which is a text file filled with instructions. The end result is a hierarchically layered binary construct. Depending on the backend, the implementation uses overlay or copy-on-write (COW) mechanisms to represent the image.

Layer example for a Tomcat application:

1. Base image (CentOS 7)
2. Install Java
3. Install Tomcat
4. Install App

The pre-built images under version control can be saved in an image registry and can then be used by the container platform.

Namespaces

Namespaces in Kubernetes represent a logical segregation of unique names for entities (Pods, Services, Deployments, ConfigMaps, etc.).

Permissions and roles can be bound on a per-namespace basis. This way, a user can control his own resources inside a namespace.

Note

Some resources are valid cluster-wise and cannot be set and controlled on a namespace basis.

Pods

A Pod is the smallest entity in Kubernetes.

It represents one instance of your running application process. The Pod consists of at least two containers, one for your application itself and another one as part of the Kubernetes design, to keep the network namespace. The so-called infrastructure container (or pause container) is therefore automatically added by

Kubernetes.

The application ports from inside the Pod are exposed via Services.

Services

A service represents a static endpoint for your application in the Pod. As a Pod and its IP address typically are considered dynamic, the IP address of the Service does not change when changing the application inside the Pod. If you scale up your Pods, you have an automatic internal load balancing towards all Pod IP addresses.

There are different kinds of Services:

- `ClusterIP` : Default virtual IP address range
- `NodePort` : Same as `ClusterIP` plus open ports on the nodes
- `LoadBalancer` : An external load balancer is created, only works in cloud environments, e.g. AWS ELB
- `ExternalName` : A DNS entry is created, also only works in cloud environments

A Service is unique inside a Namespace.

Deployment

Have a look at the [official documentation](#) .

Volume

Have a look at the [official documentation](#) .

Job

Have a look at the [official documentation](#) .

History

There is a official Kubernetes Documentary available on Youtube.

- [Kubernetes: The Documentary \[PART 1\]](#)
- [Kubernetes: The Documentary \[PART 2\]](#)

Inspired by the open source success of Docker in 2013 and seeing the need for innovation in the area of large-scale cloud computing, a handful of forward-thinking Google engineers set to work on the container orchestrator that would come to be known as Kubernetes- this new tool would forever change the way the internet is built.

These engineers overcome technical challenges, resistance to open source from within, naysayers, and intense competition from other big players in the industry.

Most engineers know about “The Container Orchestrator Wars” but most people would not be able to explain exactly what happened, and why it was Kubernetes that ultimately came out on top.

There is no topic more relevant to the current open source landscape. This film captures the story directly from the people who lived it, featuring interviews with prominent engineers from Google, Red Hat, Twitter and others.

1.1. YAML

YAML Ain't Markup Language (YAML) is a human-readable data-serialization language. YAML is not a programming language. It is mostly used for storing configuration information.

Note

Data serialization is the process of converting data objects, or object states present in complex data structures, into a stream of bytes for storage, transfer, and distribution in a form that can allow recovery of its original structure.

As you will see a lot of YAML in our Kubernetes basics course, we want to make sure you can read and write YAML. If you are not yet familiar with YAML, this introduction is waiting for you. Otherwise, feel free to skip it or come back later if you meet some less familiar YAML stuff.

This introduction is based on the [YAML Tutorial from cloudbees.com](https://cloudbees.com/yaml-tutorial/).

For more information and the full spec have a look at <https://yaml.org/>

A simple file

Let's look at a YAML file for an overview:

```
---
foo: "foo is not bar"
bar: "bar is not foo"
pi: 3.14159
awesome: true
kubernetes-birth-year: 2015
cloud-native:
  - scalable
  - dynamic
  - cloud
  - container
kubernetes:
  version: "1.22.0"
  deployed: true
  applications:
    - name: "My App"
      location: "public cloud"
```

The file starts with three dashes. These dashes indicate the start of a new YAML document. YAML supports multiple documents, and compliant parsers will recognize each set of dashes as the beginning of a new one.

Then we see the construct that makes up most of a typical YAML document: a key-value pair. `foo` is a key that points to a string value: `foo is not bar`

YAML knows four different data types:

- `foo` & `bar` are strings.
- `pi` is a floating-point number
- `awesome` is a boolean
- `kubernetes-birth-year` is an integer

You can enclose strings in single or double-quotes or no quotes at all. YAML recognizes unquoted numerals as integers or floating point.

- acend gmbh

The `cloud-native` item is an array with four elements, each denoted by an opening dash. The elements in `cloud-native` are indented with two spaces. Indentation is how YAML denotes nesting. The number of spaces can vary from file to file, but tabs are not allowed.

Finally, `kubernetes` is a dictionary that contains a string `version`, a boolean `deployed` and an array `applications` where the item of the array contains two strings.

YAML supports nesting of key-values, and mixing types.

Indentation and Whitespace

Whitespace is part of YAML's formatting. Unless otherwise indicated, newlines indicate the end of a field. You structure a YAML document with indentation. The indentation level can be one or more spaces. The specification forbids tabs because tools treat them differently.

Comments

Comments begin with a pound sign. They can appear after a document value or take up an entire line.

```
---
# This is a full line comment
foo: bar # this is a comment, too
```

YAML data types

Values in YAML's key-value pairs are scalar. They act like the scalar types in languages like Perl, Javascript, and Python. It's usually good enough to enclose strings in quotes, leave numbers unquoted, and let the parser figure it out. But that's only the tip of the iceberg. YAML is capable of a great deal more.

Key-Value Pairs and Dictionaries

The key-value is YAML's basic building block. Every item in a YAML document is a member of at least one dictionary. The key is always a string. The value is a scalar so that it can be any datatype. So, as we've already seen, the value can be a string, a number, or another dictionary.

Numeric types

YAML recognizes numeric types. We saw floating point and integers above. YAML supports several other numeric types. An integer can be decimal, hexadecimal, or octal.

```
---
foo: 12345
bar: 0x12d4
plop: 023332
```

YAML supports both fixed and exponential floating point numbers.

```
---
foo: 1230.15
bar: 12.3015e+05
```

Finally, we can represent not-a-number (NaN) or infinity.

- acend gmbh

```
---  
foo: .inf  
bar: -.Inf  
plop: .NAN
```

Foo is infinity. Bar is negative infinity, and plop is NAN.

Strings

YAML strings are Unicode. In most situations, you don't have to specify them in quotes.

```
---  
foo: this is a normal string
```

But if we want escape sequences handled, we need to use double quotes.

```
---  
foo: "this is not a normal string\n"  
bar: this is not a normal string\n
```

YAML processes the first value as ending with a carriage return and linefeed. Since the second value is not quoted, YAML treats the `\n` as two characters.

```
foo: this is not a normal string  
bar: this is not a normal string\n
```

YAML will not escape strings with single quotes, but the single quotes do avoid having string contents interpreted as document formatting. String values can span more than one line. With the fold (greater than) character, you can specify a string in a block.

```
bar: >  
  this is not a normal string it  
  spans more than  
  one line  
  see?
```

But it's interpreted without the newlines: `bar : this is not a normal string it spans more than one line see?`

The block (pipe) character has a similar function, but YAML interprets the field exactly as is.

```
---  
bar: |  
  this is not a normal string it  
  spans more than  
  one line  
  see?
```

- acend gmbh

So, we see the newlines where they are in the document.

```
bar : this is not a normal string it
      spans more than
      one line
      see?
```

Nulls

You enter nulls with a tilde or the unquoted null string literal.

```
---
foo: ~
bar: null
```

Booleans

YAML indicates boolean values with the keywords True, On and Yes for true. False is indicated with False, Off, or No.

```
---
foo: True
bar: False
light: On
TV: Off
```

Arrays

You can specify arrays or lists on a single line.

```
---
items: [ 1, 2, 3, 4, 5 ]
names: [ "one", "two", "three", "four" ]
```

Or, you can put them on multiple lines.

```
---
items:
  - 1
  - 2
  - 3
  - 4
  - 5
names:
  - "one"
  - "two"
  - "three"
  - "four"
```

- acend gmbh

The multiple line format is useful for lists that contain complex objects instead of scalars.

```
---
items:
- things:
  thing1: huey
  things2: dewey
  thing3: louie
- other things:
  key: value
```

An array can contain any valid YAML value. The values in a list do not have to be the same type.

Dictionaries

We covered dictionaries above, but there's more to them. Like arrays, you can put dictionaries inline. We saw this format above.

```
---
foo: { thing1: huey, thing2: louie, thing3: dewey }
```

We've seen them span lines before.

```
---
foo: bar
bar: foo
```

And, of course, they can be nested and hold any value.

```
---
foo:
  bar:
    - bar
    - rab
    - plop
```

2. First steps

In this lab, we will interact with the Kubernetes cluster for the first time.

Warning

Please make sure you completed *Setup* before you continue with this lab.

Login

Note

Authentication depends on the specific Kubernetes cluster environment. You may need special instructions if you are not using our lab environment. Details will be provided by your teacher.

Namespaces

Note

If you work in our acend web based environment, your Namespace has already been created and you can skip this task.

A Namespace is a logical design used in Kubernetes to organize and separate your applications, Deployments, Pods, Ingresses, Services, etc. on a top-level basis. Take a look at the [Kubernetes docs](#). Authorized users inside a namespace are able to manage those resources. Namespace names have to be unique in your cluster.

Task 2.1: Create a Namespace

Create a new namespace on the Kubernetes Cluster.. The `kubectl help` output can help you figure out the right command.

Note

Please choose an identifying name for your Namespace, e.g. your initials or name as a prefix.

We are going to use `<namespace>` as a placeholder for your created Namespace.

Solution

To create a new Namespace on your cluster use the following command:

```
kubectl create namespace <namespace>
```

Note

By using the following command, you can switch into another Namespace instead of specifying it for each `kubectl` command.

Linux:

```
kubectl config set-context $(kubectl config current-context) --namespace <namespace>
```

Windows:

```
kubectl config current-context  
SET KUBE_CONTEXT=[Insert output of the upper command]  
kubectl config set-context %KUBE_CONTEXT% --namespace <namespace>
```

Some prefer to explicitly select the Namespace for each `kubectl` command by adding `--namespace <namespace>` or `-n <namespace>`. Others prefer helper tools like `kubens` (see).

3. Deploying a container image

In this lab, we are going to deploy our first container image and look at the concepts of Pods, Services, and Deployments.

Task 3.1: Start and stop a single Pod

After we've familiarized ourselves with the platform, we are going to have a look at deploying a pre-built container image from Quay.io or any other public container registry.

First, we are going to directly start a new Pod. For this we have to define our Kubernetes Pod resource definition. Create a new file `pod_awesome-app.yaml` with the content below.

```
apiVersion: v1
kind: Pod
metadata:
  name: awesome-app
spec:
  containers:
  - image: quay.io/acend/example-web-go:latest
    imagePullPolicy: Always
    name: awesome-app
  resources:
    limits:
      cpu: 20m
      memory: 32Mi
    requests:
      cpu: 10m
      memory: 16Mi
```

Now we can apply this with:

```
kubectl apply -f pod_awesome-app.yaml --namespace <namespace>
```

The output should be:

```
pod/awesome-app created
```

Use `kubectl get pods --namespace <namespace>` in order to show the running Pod:

```
kubectl get pods --namespace <namespace>
```

Which gives you an output similar to this:

NAME	READY	STATUS	RESTARTS	AGE
awesome-app	1/1	Running	0	1m24s

Now delete the newly created Pod:

```
kubectl delete pod awesome-app --namespace <namespace>
```

Task 3.2: Create a Deployment

In some use cases it can make sense to start a single Pod. But this has its downsides and is not really a common practice. Let's look at another concept which is tightly coupled with the Pod: the so-called *Deployment*. A Deployment ensures that a Pod is monitored and checks that the number of running Pods corresponds to the number of requested Pods.

To create a new Deployment we first define our Deployment in a new file `deployment_example-web-go.yaml` with the content below.

```
apiVersion: apps/v1
kind: Deployment
metadata:
  labels:
    app: example-web-go
    name: example-web-go
spec:
  replicas: 1
  selector:
    matchLabels:
      app: example-web-go
  template:
    metadata:
      labels:
        app: example-web-go
    spec:
      containers:
        - image: quay.io/acend/example-web-go:latest
          name: example-web-go
          resources:
            requests:
              cpu: 10m
              memory: 16Mi
            limits:
              cpu: 20m
              memory: 32Mi
```

And with this we create our Deployment inside our already created namespace:

```
kubectl apply -f deployment_example-web-go.yaml --namespace <namespace>
```

The output should be:

```
deployment.apps/example-web-go created
```

We're using a simple sample application written in Go, which you can find built as an image on [Quay.io](https://quay.io/repository/acend/example-web-go) or as source code on [GitHub](https://github.com/acend/example-web-go).

Kubernetes creates the defined and necessary resources, pulls the container image (in this case from Quay.io) and deploys the Pod.

Use the command `kubectl get` with the `-w` parameter in order to get the requested resources and afterward watch for changes.

Note

The `kubectl get -w` command will never end unless you terminate it with `CTRL-C`.

```
kubectl get pods -w --namespace <namespace>
```

Note

Instead of using the `-w` parameter you can also use the `watch` command which should be available on most Linux distributions:

```
watch kubectl get pods --namespace <namespace>
```

This process can last for some time depending on your internet connection and if the image is already available locally.

Note

If you want to create your own container images and use them with Kubernetes, you definitely should have a look at [these best practices](#) and apply them. This image creation guide may be for OpenShift, however it also applies to Kubernetes and other container platforms.

Creating Kubernetes resources

There are two fundamentally different ways to create Kubernetes resources. You've already seen one way: Writing the resource's definition in YAML (or JSON) and then applying it on the cluster using `kubectl apply`.

The other variant is to use helper commands. These are more straightforward: You don't have to copy a YAML definition from somewhere else and then adapt it. However, the result is the same. The helper commands just simplify the process of creating the YAML definitions.

As an example, let's look at creating above deployment, this time using a helper command instead. If you already created the Deployment using above YAML definition, you don't have to execute this command:

```
kubectl create deployment example-web-go --image=quay.io/acend/example-web-go:latest --namespace <namespace>
```

It's important to know that these helper commands exist. However, in a world where GitOps concepts have an ever-increasing presence, the idea is not to constantly create these resources with helper commands. Instead, we save the resources' YAML definitions in a Git repository and leave the creation and management of those resources to a tool.

Task 3.3: Viewing the created resources

Display the created Deployment using the following command:

```
kubectl get deployments --namespace <namespace>
```

- acend gmbh

A [Deployment](#) defines the following facts:

- Update strategy: How application updates should be executed and how the Pods are exchanged
- Containers
 - Which image should be deployed
 - Environment configuration for Pods
 - ImagePullPolicy
- The number of Pods/Replicas that should be deployed

By using the `-o` (or `--output`) parameter we get a lot more information about the deployment itself. You can choose between YAML and JSON formatting by indicating `-o yaml` or `-o json`. In this training we are going to use YAML, but please feel free to replace `yaml` with `json` if you prefer.

```
kubectl get deployment example-web-go -o yaml --namespace <namespace>
```

After the image has been pulled, Kubernetes deploys a Pod according to the Deployment:

```
kubectl get pods --namespace <namespace>
```

which gives you an output similar to this:

NAME	READY	STATUS	RESTARTS	AGE
example-web-go-69b658f647-xnm94	1/1	Running	0	39s

The Deployment defines that one replica should be deployed — which is running as we can see in the output. This Pod is not yet reachable from outside the cluster.

4. Exposing a service

In this lab, we are going to make the freshly deployed application from the last lab available online.

Task 4.1: Create a ClusterIP Service

The command `kubectl apply -f deployment_example-web-go.yaml` from the last lab creates a Deployment but no Service. A Kubernetes Service is an abstract way to expose an application running on a set of Pods as a network service. For some parts of your application (for example, frontends) you may want to expose a Service to an external IP address which is outside your cluster.

Kubernetes `ServiceTypes` allow you to specify what kind of Service you want. The default is `ClusterIP`.

Type values and their behaviors are:

- `ClusterIP` : Exposes the Service on a cluster-internal IP. Choosing this value only makes the Service reachable from within the cluster. This is the default `ServiceType`.
- `NodePort` : Exposes the Service on each Node's IP at a static port (the `NodePort`). A `ClusterIP` Service, to which the `NodePort` Service routes, is automatically created. You'll be able to contact the `NodePort` Service from outside the cluster, by requesting `<NodeIP>:<NodePort>`.
- `LoadBalancer` : Exposes the Service externally using a cloud provider's load balancer. `NodePort` and `ClusterIP` Services, to which the external load balancer routes, are automatically created.
- `ExternalName` : Maps the Service to the contents of the `externalName` field (e.g. `foo.bar.example.com`), by returning a CNAME record with its value. No proxying of any kind is set up.

You can also use Ingress to expose your Service. Ingress is not a Service type, but it acts as the entry point for your cluster. [Ingress](#) exposes HTTP and HTTPS routes from outside the cluster to services within the cluster. Traffic routing is controlled by rules defined on the Ingress resource. An Ingress may be configured to give Services externally reachable URLs, load balance traffic, terminate SSL / TLS, and offer name-based virtual hosting. An Ingress controller is responsible for fulfilling the route, usually with a load balancer, though it may also configure your edge router or additional frontends to help handle the traffic.

In order to create an Ingress, we first need to create a Service of type [ClusterIP](#).

To create the Service add a new file `svc-web-go.yaml` with the following content:

```
---
apiVersion: v1
kind: Service
metadata:
  labels:
    app: example-web-go
    name: example-web-go
spec:
  ports:
    - port: 5000
      protocol: TCP
      targetPort: 5000
  selector:
    app: example-web-go
  type: ClusterIP
```

And then apply the file with:

- acend gmbh

```
kubectl apply -f svc-web-go.yaml --namespace <namespace>
```

There is also an imperative command to create a service and expose your application which can be used instead of the yaml file with the `kubectl apply ...` command

```
kubectl expose deployment example-web-go --type=ClusterIP --name=example-web-go --port=5000 --target-port=5000 --namespace <namespace>
```

Let's have a more detailed look at our Service:

```
kubectl get services --namespace <namespace>
```

Which gives you an output similar to this:

NAME	TYPE	CLUSTER-IP	EXTERNAL-IP	PORT(S)	AGE
example-web-go	ClusterIP	10.43.91.62	<none>	5000/TCP	

Note

Service IP (CLUSTER-IP) addresses stay the same for the duration of the Service's lifespan.

By executing the following command:

```
kubectl get service example-web-go -o yaml --namespace <namespace>
```

You get additional information:

```
apiVersion: v1
kind: Service
metadata:
  ...
  labels:
    app: example-web-go
  managedFields:
    ...
  name: example-web-go
  namespace: <namespace>
  ...
spec:
  clusterIP: 10.43.91.62
  externalTrafficPolicy: Cluster
  ports:
    - port: 5000
      protocol: TCP
      targetPort: 5000
  selector:
    app: example-web-go
  sessionAffinity: None
  type: ClusterIP
status:
  loadBalancer: {}
```

- acend gmbh

The Service's `selector` defines which Pods are being used as Endpoints. This happens based on labels. Look at the configuration of Service and Pod in order to find out what maps to what:

```
kubectl get service example-web-go -o yaml --namespace <namespace>
```

```
...
  selector:
    app: example-web-go
...
```

With the following command you get details from the Pod:

Note

First, get all Pod names from your namespace with (`kubectl get pods --namespace <namespace>`) and then replace `<pod>` in the following command. If you have installed and configured the bash completion, you can also press the TAB key for autocompletion of the Pod's name.

```
kubectl get pod <pod> -o yaml --namespace <namespace>
```

Let's have a look at the label section of the Pod and verify that the Service selector matches the Pod's labels:

```
...
  labels:
    app: example-web-go
...
```

This link between Service and Pod can also be displayed in an easier fashion with the `kubectl describe` command:

```
kubectl describe service example-web-go --namespace <namespace>
```

```
Name:                example-web-go
Namespace:            example-ns
Labels:               app=example-web-go
Annotations:          <none>
Selector:             app=example-web-go
Type:                 ClusterIP
IP:                  10.39.240.212
Port:                 <unset> 5000/TCP
TargetPort:           5000/TCP
Endpoints:            10.36.0.8:5000
Session Affinity:     None
External Traffic Policy: Cluster
Events:
  Type     Reason          Age   From          Message
  ----     -
  ...
```

The `Endpoints` show the IP addresses of all currently matched Pods.

Task 4.2: Expose the Service

With the ClusterIP Service ready, we can now create the Ingress resource.

In order to create the Ingress resource, we first need to create the file `ing-example-web-go.yaml` and change the `host` entry to match your environment:

```
---
apiVersion: networking.k8s.io/v1
kind: Ingress
metadata:
  name: example-web-go
spec:
  rules:
    - host: example-web-go-<namespace>.<appdomain>
      http:
        paths:
          - path: /
            pathType: Prefix
            backend:
              service:
                name: example-web-go
                port:
                  number: 5000
  tls:
    - hosts:
        - example-web-go-<namespace>.<appdomain>
```

As you see in the resource definition at `spec.rules[0].http.paths[0].backend.service.name` we use the previously created `example-web-go` ClusterIP Service.

Let's create the Ingress resource with:

```
kubectl apply -f ing-example-web-go.yaml --namespace <namespace>
```

Afterwards, we are able to access our freshly created Ingress at `http://example-web-go-<namespace>.<appdomain>`

Task 4.3: Expose as NodePort

Note

This is an advanced lab, so feel free to skip this. NodePorts are usually not used for http-based applications as we use the layer 7-based Ingress resource. Only for non-http based applications, a NodePort might be a suitable alternative.

There's a second option to make a Service accessible from outside: Use a [NodePort](#) .

In order to switch the Service type, change the existing `ClusterIP` Service by updating our Service definition in file `svc-web-go.yaml` to:

- acend gmbh

```
apiVersion: v1
kind: Service
metadata:
  labels:
    app: example-web-go
    name: example-web-go
spec:
  ports:
    - port: 5000
      protocol: TCP
      targetPort: 5000
  selector:
    app: example-web-go
  type: NodePort
```

And then apply again with:

```
kubectl apply -f svc-web-go.yaml --namespace <namespace>
```

Let's have a more detailed look at our new `NodePort` Service:

```
kubectl get services --namespace <namespace>
```

Which gives you an output similar to this:

NAME	TYPE	CLUSTER-IP	EXTERNAL-IP	PORT(S)	AGE
example-web-go	NodePort	10.43.91.62	<none>	5000:30692/TCP	

The `NodePort` number is assigned by Kubernetes and stays the same as long as the Service is not deleted. A `NodePort` Service is more suitable for infrastructure tools than for public URLs.

Open `http://<node-ip>:<node-port>` in your browser or use `curl http://<node-ip>:<node-port>` when the public ip is not available in your browser. You can use any node IP as the Service is exposed on all nodes using the same `NodePort`. Use `kubectl get nodes -o wide` to display the IPs (`INTERNAL-IP` or `EXTERNAL-IP`) of the available nodes. Depending on your environment, use the internal or external (public) ip address.

```
kubectl get node -o wide
```

The output may vary depending on your setup:

NAME	STATUS	ROLES	AGE	VERSION	INTERNAL-IP	EXTERNAL-IP	OS-IMAGE	KE
RHEL-VERSION		CONTAINER-RUNTIME						
lab-1	Ready	controlplane,etcd,worker	150m	v1.17.4	5.102.145.142	<none>	Ubuntu 18.04.3 LTS	4.15.0-
66-generic		docker://19.3.8						
lab-2	Ready	controlplane,etcd,worker	150m	v1.17.4	5.102.145.77	<none>	Ubuntu 18.04.3 LTS	4.15.0-
66-generic		docker://19.3.8						
lab-3	Ready	controlplane,etcd,worker	150m	v1.17.4	5.102.145.148	<none>	Ubuntu 18.04.3 LTS	4.15.0-
66-generic		docker://19.3.8						

Task 4.4: For fast learners

Have a closer look at the resources created in your namespace `<namespace>` with the following commands and try to understand them:

```
kubectl describe namespace <namespace>
```

```
kubectl get all --namespace <namespace>
```

```
kubectl describe <resource> <name> --namespace <namespace>
```

```
kubectl get <resource> <name> -o yaml --namespace <namespace>
```

5. Scaling

In this lab, we are going to show you how to scale applications on Kubernetes. Furthermore, we show you how Kubernetes makes sure that the number of requested Pods is up and running and how an application can tell the platform that it is ready to receive requests.

Note

This lab does not depend on previous labs. You can start with an empty Namespace.

Task 5.1: Scale the example application

Create a new Deployment in your Namespace. So again, let's define the Deployment using YAML in a file `deployment_example-web-app.yaml` with the following content:

```
apiVersion: apps/v1
kind: Deployment
metadata:
  labels:
    app: example-web-app
    name: example-web-app
spec:
  replicas: 1
  selector:
    matchLabels:
      app: example-web-app
  strategy:
    rollingUpdate:
      maxSurge: 25%
      maxUnavailable: 0
    type: RollingUpdate
  template:
    metadata:
      labels:
        app: example-web-app
    spec:
      containers:
        - image: quay.io/acend/example-web-python:latest
          name: example-web-app
          resources:
            limits:
              cpu: 100m
              memory: 128Mi
            requests:
              cpu: 50m
              memory: 128Mi
```

and then apply with:

```
kubectl apply -f deployment_example-web-app.yaml --namespace <namespace>
```

If we want to scale our example application, we have to tell the Deployment that we want to have three running replicas instead of one. Let's have a closer look at the existing ReplicaSet:

```
kubectl get replicaset --namespace <namespace>
```

- acend gmbh

Which will give you an output similar to this:

NAME	DESIRED	CURRENT	READY	AGE
example-web-app-86d9d584f8	1	1	1	110s

Or for even more details:

```
kubectl get replicaset <replicaset> -o yaml --namespace <namespace>
```

The ReplicaSet shows how many instances of a Pod are desired, current and ready.

Now we scale our application to three replicas:

```
kubectl scale deployment example-web-app --replicas=3 --namespace <namespace>
```

Check the number of desired, current and ready replicas:

```
kubectl get replicaset --namespace <namespace>
```

NAME	DESIRED	CURRENT	READY	AGE
example-web-app-86d9d584f8	3	3	3	4m33s

Look at how many Pods there are:

```
kubectl get pods --namespace <namespace>
```

Which gives you an output similar to this:

NAME	READY	STATUS	RESTARTS	AGE
example-web-app-86d9d584f8-7vjcj	1/1	Running	0	5m2s
example-web-app-86d9d584f8-hbvlv	1/1	Running	0	31s
example-web-app-86d9d584f8-qg499	1/1	Running	0	31s

Note

Kubernetes even supports [autoscaling](#) .

As we changed the number of replicas with the `kubectl scale deployment` command, the `example-web-app` Deployment now differs from your local `deployment_example-web-app.yaml` file. Change your local `deployment_example-web-app.yaml` file to match the current number of replicas and update the value `replicas` to `3` :

- acend gmbh

```
[...]
metadata:
  labels:
    app: example-web-app
  name: example-web-app
spec:
  replicas: 3
  selector:
    matchLabels:
      app: example-web-app
[...]
```

Check for uninterruptible Deployments

Now we create a new Service of the type `ClusterIP`. Create a new file `svc-example-app.yaml` with the following content:

```
---
apiVersion: v1
kind: Service
metadata:
  labels:
    app: example-web-app
  name: example-web-app
spec:
  ports:
    - port: 5000
      protocol: TCP
      targetPort: 5000
  selector:
    app: example-web-app
  type: ClusterIP
```

and apply the file with:

```
kubectl apply -f svc-example-app.yaml --namespace <namespace>
```

Then we add the Ingress to access our application. Create a new file `ing-example-web-app.yaml` with the following content:

- acend gmbh

```
---
apiVersion: networking.k8s.io/v1
kind: Ingress
metadata:
  name: example-web-app
spec:
  rules:
    - host: example-web-app-<namespace>.<appdomain>
      http:
        paths:
          - path: /
            pathType: Prefix
            backend:
              service:
                name: example-web-app
                port:
                  number: 5000
  tls:
    - hosts:
        - example-web-app-<namespace>.<appdomain>
```

Apply this Ingress definition using, e.g.:

```
kubectl apply -f ing-example-web-app.yaml --namespace <namespace>
```

Let's look at our Service. We should see all three corresponding Endpoints:

```
kubectl describe service example-web-app --namespace <namespace>
```

```
Name:                example-web-app
Namespace:           acend-scale
Labels:              app=example-web-app
Annotations:         <none>
Selector:            app=example-web-app
Type:               ClusterIP
IP:                 10.39.245.205
Port:               <unset> 5000/TCP
TargetPort:         5000/TCP
Endpoints:          10.36.0.10:5000,10.36.0.11:5000,10.36.0.9:5000
Session Affinity:   None
External Traffic Policy: Cluster
Events:
  Type    Reason          Age   From          Message
  ----    -
  ----    -
```

Scaling of Pods is fast as Kubernetes simply creates new containers.

You can check the availability of your Service while you scale the number of replicas up and down in your browser: <http://example-web-app-<namespace>.<appdomain>> .

Now, execute the corresponding loop command for your operating system in another console.

Linux:

```
URL=$(kubectl get ingress example-web-app -o go-template="{{ (index .spec.rules 0).host }}" --namespace <namespace>)
while true; do sleep 1; curl -s https://${URL}/pod/; date "+ TIME: %H:%M:%S,%3N"; done
```

Windows PowerShell:

```
while(1) {  
    Start-Sleep -s 1  
    Invoke-RestMethod http://<URL>/pod/  
    Get-Date -Uformat "+ TIME: %H:%M:%S,%3N"  
}
```

Scale from 3 replicas to 1. The output shows which Pod is still alive and is responding to requests:

```
example-web-app-86d9d584f8-7vjcj TIME: 17:33:07,289  
example-web-app-86d9d584f8-7vjcj TIME: 17:33:08,357  
example-web-app-86d9d584f8-hbvlv TIME: 17:33:09,423  
example-web-app-86d9d584f8-7vjcj TIME: 17:33:10,494  
example-web-app-86d9d584f8-qg499 TIME: 17:33:11,559  
example-web-app-86d9d584f8-hbvlv TIME: 17:33:12,629  
example-web-app-86d9d584f8-qg499 TIME: 17:33:13,695  
example-web-app-86d9d584f8-hbvlv TIME: 17:33:14,771  
example-web-app-86d9d584f8-hbvlv TIME: 17:33:15,840  
example-web-app-86d9d584f8-7vjcj TIME: 17:33:16,912  
example-web-app-86d9d584f8-7vjcj TIME: 17:33:17,980  
example-web-app-86d9d584f8-7vjcj TIME: 17:33:19,051  
example-web-app-86d9d584f8-7vjcj TIME: 17:33:20,119  
example-web-app-86d9d584f8-7vjcj TIME: 17:33:21,182  
example-web-app-86d9d584f8-7vjcj TIME: 17:33:22,248  
example-web-app-86d9d584f8-7vjcj TIME: 17:33:23,313  
example-web-app-86d9d584f8-7vjcj TIME: 17:33:24,377  
example-web-app-86d9d584f8-7vjcj TIME: 17:33:25,445  
example-web-app-86d9d584f8-7vjcj TIME: 17:33:26,513
```

The requests get distributed amongst the three Pods. As soon as you scale down to one Pod, there should be only one remaining Pod that responds.

Let's make another test: What happens if you start a new Deployment while our request generator is still running?

```
kubectl rollout restart deployment example-web-app --namespace <namespace>
```

During a short period we won't get a response:

```
example-web-app-86d9d584f8-7vjcj TIME: 17:37:24,121  
example-web-app-86d9d584f8-7vjcj TIME: 17:37:25,189  
example-web-app-86d9d584f8-7vjcj TIME: 17:37:26,262  
example-web-app-86d9d584f8-7vjcj TIME: 17:37:27,328  
example-web-app-86d9d584f8-7vjcj TIME: 17:37:28,395  
example-web-app-86d9d584f8-7vjcj TIME: 17:37:29,459  
example-web-app-86d9d584f8-7vjcj TIME: 17:37:30,531  
example-web-app-86d9d584f8-7vjcj TIME: 17:37:31,596  
example-web-app-86d9d584f8-7vjcj TIME: 17:37:32,662  
# no answer  
example-web-app-f4c5dd8fc-4nx2t TIME: 17:37:33,729  
example-web-app-f4c5dd8fc-4nx2t TIME: 17:37:34,794  
example-web-app-f4c5dd8fc-4nx2t TIME: 17:37:35,862  
example-web-app-f4c5dd8fc-4nx2t TIME: 17:37:36,929  
example-web-app-f4c5dd8fc-4nx2t TIME: 17:37:37,995  
example-web-app-f4c5dd8fc-4nx2t TIME: 17:37:39,060  
example-web-app-f4c5dd8fc-4nx2t TIME: 17:37:40,118  
example-web-app-f4c5dd8fc-4nx2t TIME: 17:37:41,187
```

In our example, we use a very lightweight Pod. If we had used a more heavyweight Pod that needed a

- acend gmbh

longer time to respond to requests, we would of course see a larger gap. An example for this would be a Java application with a startup time of 30 seconds:

```
example-spring-boot-2-73aln TIME: 16:48:25,251
example-spring-boot-2-73aln TIME: 16:48:26,305
example-spring-boot-2-73aln TIME: 16:48:27,400
example-spring-boot-2-73aln TIME: 16:48:28,463
example-spring-boot-2-73aln TIME: 16:48:29,507
<html><body><h1>503 Service Unavailable</h1>
No server is available to handle this request.
</body></html>
TIME: 16:48:33,562
<html><body><h1>503 Service Unavailable</h1>
No server is available to handle this request.
</body></html>
TIME: 16:48:34,601
...
example-spring-boot-3-tjdkj TIME: 16:49:20,114
example-spring-boot-3-tjdkj TIME: 16:49:21,181
example-spring-boot-3-tjdkj TIME: 16:49:22,231
```

It is even possible that the Service gets down, and the routing layer responds with the status code 503 as can be seen in the example output above.

In the following chapter we are going to look at how a Service can be configured to be highly available.

Uninterruptible Deployments

The [rolling update strategy](#) makes it possible to deploy Pods without interruption. The rolling update strategy means that the new version of an application gets deployed and started. As soon as the application says it is ready, Kubernetes forwards requests to the new instead of the old version of the Pod, and the old Pod gets terminated.

Additionally, [container health checks](#) help Kubernetes to precisely determine what state the application is in.

Basically, there are two different kinds of checks that can be implemented:

- Liveness probes are used to find out if an application is still running
- Readiness probes tell us if the application is ready to receive requests (which is especially relevant for the above-mentioned rolling updates)

These probes can be implemented as HTTP checks, container execution checks (the execution of a command or script inside a container) or TCP socket checks.

In our example, we want the application to tell Kubernetes that it is ready for requests with an appropriate readiness probe.

Our example application has a health check context named health: `http://<node-ip>:<node-port>/health`

Task 5.2: Availability during deployment

In our deployment configuration inside the rolling update strategy section, we define that our application always has to be available during an update: `maxUnavailable: 0`

Now insert the readiness probe at `.spec.template.spec.containers` above the `resources` line in your local `deployment_example-web-app.yaml` File:

- acend gmbh

```
...
containers:
- image: quay.io/acend/example-web-python:latest
  imagePullPolicy: Always
  name: example-web-app
  # start to copy here
  readinessProbe:
    httpGet:
      path: /health
      port: 5000
      scheme: HTTP
    initialDelaySeconds: 10
    timeoutSeconds: 1
  # stop to copy here
  resources:
    limits:
      cpu: 100m
      memory: 128Mi
    requests:
      cpu: 50m
      memory: 128Mi
...
```

apply the file with:

```
kubectl apply -f deployment_example-web-app.yaml --namespace <namespace>
```

We are now going to verify that a redeployment of the application does not lead to an interruption.

Set up the loop again to periodically check the application's response (you don't have to set the `$URL` variable again if it is still defined):

```
URL=$(kubectl get ingress example-web-app -o go-template="{{ (index .spec.rules 0).host }}" --namespace <namespace>)
while true; do sleep 1; curl -s http://$URL/pod/; date "+ TIME: %H:%M:%S,%3N"; done
```

Windows PowerShell:

```
while(1) {
    Start-Sleep -s 1
    Invoke-RestMethod http://[URL]/pod/
    Get-Date -Uformat "+ TIME: %H:%M:%S,%3N"
}
```

Restart your Deployment with:

```
kubectl rollout restart deployment example-web-app --namespace <namespace>
```

Self-healing

Via the Replicaset we told Kubernetes how many replicas we want. So what happens if we simply delete a

- acend gmbh

Pod?

Look for a running Pod (status `RUNNING`) that you can bear to kill via `kubectl get pods` .

Show all Pods and watch for changes:

```
kubectl get pods -w --namespace <namespace>
```

Now delete a Pod (in another terminal) with the following command:

```
kubectl delete pod <pod> --namespace <namespace>
```

Observe how Kubernetes instantly creates a new Pod in order to fulfill the desired number of running instances.

6. Troubleshooting

This lab helps you troubleshoot your application and shows you some tools to make troubleshooting easier.

Logging into a container

Running containers should be treated as immutable infrastructure and should therefore not be modified. However, there are some use cases in which you have to log into your running container. Debugging and analyzing is one example for this.

Task 6.1: Shell into Pod

With Kubernetes you can open a remote shell into a Pod without installing SSH by using the command `kubectl exec`. The command can also be used to execute any command in a Pod. If you want to get a shell to a running container, you will additionally need the parameters `-it`. These set up an interactive session where you can supply input to the process inside the container.

Note

If you're using Git Bash on Windows, you need to append the command with `winpty`.

Choose a Pod with `kubectl get pods --namespace <namespace>` and execute the following command:

```
kubectl exec -it <pod> --namespace <namespace> -- /bin/bash
```

Note

If Bash is not available in the Pod you can fallback to `- sh` instead of `- /bin/bash`.

You now have a running shell session inside the container in which you can execute every binary available, e.g.:

```
ls -l
```

```
total 12
-rw-r--r-- 1 10020700 root      8192 Nov 27 15:12 helloos.db
-rwxrwsr-x 1 web      root      2454 Oct  5 08:55 run.py
drwxrwsr-x 1 web      root        17 Oct  5 08:55 static
drwxrwsr-x 1 web      root        63 Oct  5 08:55 templates
```

With `exit` or `CTRL+d` you can leave the container and close the connection:

```
exit
```

Task 6.2: Single commands

Single commands inside a container can also be executed with `kubectl exec` :

```
kubectl exec <pod> --namespace <namespace> -- env
```

Example:

```
$ kubectl exec example-web-app-69b658f647-xnm94 --namespace <namespace> -- env
PATH=/usr/local/sbin:/usr/local/bin:/usr/sbin:/usr/bin:/sbin:/bin
HOSTNAME=example-web-app-xnm94
KUBERNETES_SERVICE_PORT_DNS_TCP=53
KUBERNETES_PORT_443_TCP_PROTO=tcp
KUBERNETES_PORT_443_TCP_ADDR=172.30.0.1
KUBERNETES_PORT_53_UDP_PROTO=udp
KUBERNETES_PORT_53_TCP=tcp://172.30.0.1:53
...
```

Watching log files

Log files of a Pod can be shown with the following command:

```
kubectl logs <pod> --namespace <namespace>
```

The parameter `-f` allows you to follow the log file (same as `tail -f`). With this, log files are streamed and new entries are shown immediately.

When a Pod is in state `CrashLoopBackOff` it means that although multiple attempts have been made, no container inside the Pod could be started successfully. Now even though no container might be running at the moment the `kubectl logs` command is executed, there is a way to view the logs the application might have generated. This is achieved using the `-p` or `--previous` parameter.

Note

This command will only work on pods that had container restarts. You can check the `RESTARTS` column in the `kubectl get pods` output if this is the case.

```
kubectl logs -p <pod> --namespace <namespace>
```

Task 6.3: Port forwarding

Kubernetes allows you to forward arbitrary ports to your development workstation. This allows you to access admin consoles, databases, etc., even when they are not exposed externally. Port forwarding is handled by the Kubernetes control plane nodes and therefore tunneled from the client via HTTPS. This allows you to access the Kubernetes platform even when there are restrictive firewalls or proxies between your workstation and Kubernetes.

- acend gmbh

Get the name of the Pod:

```
kubectl get pod --namespace <namespace>
```

Then execute the port forwarding command using the Pod's name:

Note

Best run this command in a separate shell, or in the background by adding a "&" at the end of the command.

```
kubectl port-forward <pod> 5000:5000 --namespace <namespace>
```

Don't forget to change the Pod name to your own installation. If configured, you can use auto-completion.

The output of the command should look like this:

```
Forwarding from 127.0.0.1:5000 -> 5000
Forwarding from [::1]:5000 -> 5000
```

Note

Use the additional parameter `--address <IP address>` (where `<IP address>` refers to a NIC's IP address from your local workstation) if you want to access the forwarded port from outside your own local workstation.

The application is now available with the following link: <http://localhost:5000/> . Or try a `curl` command:

```
curl localhost:5000
```

With the same concept you can access databases from your local workstation or connect your local development environment via remote debugging to your application in the Pod.

[This documentation page](#) offers some more details about port forwarding.

Note

The `kubectl port-forward` process runs as long as it is not terminated by the user. So when done, stop it with `CTRL-C`.

Events

Kubernetes maintains an event log with high-level information on what's going on in the cluster. It's possible that everything looks okay at first but somehow something seems stuck. Make sure to have a look at the events because they can give you more information if something is not working as expected.

Use the following command to list the events in chronological order:

- acend gmbh

```
kubectl get events --sort-by=.metadata.creationTimestamp --namespace <namespace>
```

Dry-run

To help verify changes, you can use the optional `kubectl` flag `--dry-run=client -o yaml` to see the rendered YAML definition of your Kubernetes objects, without sending it to the API.

The following `kubectl` subcommands support this flag (non-final list):

- `apply`
- `create`
- `expose`
- `patch`
- `replace`
- `run`
- `set`

For example, we can use the `--dry-run=client` flag to create a template for our Deployment:

```
kubectl create deployment example-web-app --image=quay.io/acend/example-web-python:latest --namespace acend-test --dry-run=client -o yaml
```

The result is the following YAML output:

```
apiVersion: apps/v1
kind: Deployment
metadata:
  creationTimestamp: null
  labels:
    app: example-web-app
  name: example-web-app
  namespace: acend-test
spec:
  replicas: 1
  selector:
    matchLabels:
      app: example-web-app
  strategy: {}
  template:
    metadata:
      creationTimestamp: null
      labels:
        app: example-web-app
    spec:
      containers:
        - image: quay.io/acend/example-web-python:latest
          name: example-web
          resources: {}
status: {}
```

kubectl API requests

If you want to see the HTTP requests `kubectl` sends to the Kubernetes API in detail, you can use the optional flag `--v=10`.

For example, to see the API request for creating a deployment:

The resulting output looks like this:

As you can see, the output conveniently contains the corresponding `curl` commands which we could use in our own code, tools, pipelines etc.

If you created the deployment to see the output, you can delete it again as it's not used anywhere else (which is also the reason why the replicas are set to 0):

38 / 103

Progress

At this point, you are able to visualize your progress on the labs by browsing through the following page <http://localhost:5000/progress>

If you are not able to open your awesome-app with localhost, because you are using a webshell, you can also use the ingress address: `https://example-web-app-<namespace>.<appdomain>/progress` to access the dashboard.

You may need to set some extra permissions to let the dashboard monitor your progress. Have fun!

```
kubectl create rolebinding progress --clusterrole=view --serviceaccount=<namespace>:default --namespace=<namespace>
```

7. Attaching a database

Numerous applications are stateful in some way and want to save data persistently, be it in a database, as files on a filesystem or in an object store. In this lab, we are going to create a MariaDB database and configure our application to store its data in it.

Task 7.1: Instantiate a MariaDB database

We are first going to create a so-called *Secret* in which we store sensitive data. The secret will be used to access the database and also to create the initial database.

```
kubectl create secret generic mariadb \
--from-literal=database-name=acend_exempledb \
--from-literal=database-password=mysqlpassword \
--from-literal=database-root-password=mysqlrootpassword \
--from-literal=database-user=acend_user \
--namespace <namespace>
```

The Secret contains the database name, user, password, and the root password. However, these values will neither be shown with `kubectl get` nor with `kubectl describe` :

```
kubectl get secret mariadb --output yaml --namespace <namespace>
```

```
apiVersion: v1
data:
  database-name: YWN1bmQtZXhhbXBsZS1kYg==
  database-password: bXlzcWxwYXNzd29yZA==
  database-root-password: bXlzcWxyb290cGFzc3dvcmQ=
  database-user: YWN1bmRfdXNlcg==
kind: Secret
metadata:
  ...
type: Opaque
```

The reason is that all the values in the `.data` section are base64 encoded. Even though we cannot see the true values, they can easily be decoded:

```
echo "YWN1bmQtZXhhbXBsZS1kYg==" | base64 -d
```

Note

By default, Secrets are not encrypted!

However, both [OpenShift](#) and [Kubernetes \(1.13 and later\)](#) offer the capability to encrypt data in etcd.

Another option would be the use of a secrets management solution like [Vault by HashiCorp](#) .

We are now going to create a Deployment and a Service. As a first example, we use a database without persistent storage. Only use an ephemeral database for testing purposes as a restart of the Pod leads to data loss. We are going to look at how to persist this data in a persistent volume later on.

- acend gmbh

As we had seen in the earlier labs, all resources like Deployments, Services, Secrets and so on can be displayed in YAML or JSON format. It doesn't end there, capabilities also include the creation and exportation of resources using YAML or JSON files.

In our case we want to create a Deployment and Service for our MariaDB database. Save this snippet as `mariadb.yaml` :

- acend gmbh

```
---
apiVersion: v1
kind: Service
metadata:
  name: mariadb
  labels:
    app: mariadb
spec:
  ports:
    - port: 3306
  selector:
    app: mariadb
---
apiVersion: apps/v1
kind: Deployment
metadata:
  name: mariadb
  labels:
    app: mariadb
spec:
  selector:
    matchLabels:
      app: mariadb
  strategy:
    type: Recreate
  template:
    metadata:
      labels:
        app: mariadb
    spec:
      containers:
        - image: mariadb:10.5
          name: mariadb
          args:
            - "--ignore-db-dir=lost+found"
          env:
            - name: MYSQL_USER
              valueFrom:
                secretKeyRef:
                  key: database-user
                  name: mariadb
            - name: MYSQL_PASSWORD
              valueFrom:
                secretKeyRef:
                  key: database-password
                  name: mariadb
            - name: MYSQL_ROOT_PASSWORD
              valueFrom:
                secretKeyRef:
                  key: database-root-password
                  name: mariadb
            - name: MYSQL_DATABASE
              valueFrom:
                secretKeyRef:
                  key: database-name
                  name: mariadb
          livenessProbe:
            tcpSocket:
              port: 3306
          ports:
            - containerPort: 3306
              name: mariadb
      resources:
        limits:
          cpu: 500m
          memory: 512Mi
        requests:
          cpu: 50m
          memory: 128Mi
```

Apply it with:

- acend gmbh

```
kubectl apply -f mariadb.yaml --namespace <namespace>
```

As soon as the container image for `mariadb:10.5` has been pulled, you will see a new Pod using `kubectl get pods`.

The environment variables defined in the deployment configure the MariaDB Pod and how our frontend will be able to access it.

The interesting thing about Secrets is that they can be reused, e.g., in different Deployments. We could extract all the plaintext values from the Secret and put them as environment variables into the Deployments, but it's way easier to instead simply refer to its values inside the Deployment (as in this lab) like this:

```
...
spec:
  template:
    spec:
      containers:
      - name: mariadb
        env:
        - name: MYSQL_USER
          valueFrom:
            secretKeyRef:
              key: database-user
              name: mariadb
        - name: MYSQL_PASSWORD
          valueFrom:
            secretKeyRef:
              key: database-password
              name: mariadb
        - name: MYSQL_ROOT_PASSWORD
          valueFrom:
            secretKeyRef:
              key: database-root-password
              name: mariadb
        - name: MYSQL_DATABASE
          valueFrom:
            secretKeyRef:
              key: database-name
              name: mariadb
...

```

Above lines are an excerpt of the MariaDB Deployment. Most parts have been cut out to focus on the relevant lines: The references to the `mariadb` Secret. As you can see, instead of directly defining environment variables you can refer to a specific key inside a Secret. We are going to make further use of this concept for our Python application.

Task 7.3: Attach the database to the application

By default, our `example-web-app` application uses an SQLite memory database.

However, this can be changed by defining the following environment variable to use the newly created MariaDB database:

```
#MYSQL_URI=mysql://<user>:<password>@<host>/<database>
MYSQL_URI=mysql://acend_user:mysqlpassword@mariadb/acend_exampleddb
```

The connection string our `example-web-app` application uses to connect to our new MariaDB, is a concatenated string from the values of the `mariadb` Secret.

- acend gmbh

For the actual MariaDB host, you can either use the MariaDB Service's ClusterIP or DNS name as the address. All Services and Pods can be resolved by DNS using their name.

The following commands set the environment variables for the deployment configuration of the `example-web-app` application:

Warning

Depending on the shell you use, the following `set env` command works but inserts too many apostrophes! Check the deployment's environment variable afterwards or directly edit it as described further down below.

```
kubectl set env --from=secret/mariadb --prefix=MYSQL_ deploy/example-web-app --namespace <namespace>
```

and

```
kubectl set env deploy/example-web-app MYSQL_URI='mysql://$(MYSQL_DATABASE_USER):$(MYSQL_DATABASE_PASSWORD)@mariadb/$(MYSQL_DATABASE_NAME)' --namespace <namespace>
```

The first command inserts the values from the Secret, the second finally uses these values to put them in the environment variable `MYSQL_URI` which the application considers.

You can also do the changes by directly editing your local `deployment_example-web-app.yaml` file. Find the section which defines the containers. You should find it under:

```
...
spec:
...
  template:
...
    spec:
      containers:
      - image: ...
...

```

The dash before `image:` defines the beginning of a new container definition. The following specifications should be inserted into this container definition:

- acend gmbh

```
env:
- name: MYSQL_DATABASE_NAME
  valueFrom:
    secretKeyRef:
      key: database-name
      name: mariadb
- name: MYSQL_DATABASE_PASSWORD
  valueFrom:
    secretKeyRef:
      key: database-password
      name: mariadb
- name: MYSQL_DATABASE_ROOT_PASSWORD
  valueFrom:
    secretKeyRef:
      key: database-root-password
      name: mariadb
- name: MYSQL_DATABASE_USER
  valueFrom:
    secretKeyRef:
      key: database-user
      name: mariadb
- name: MYSQL_URI
  value: mysql://$(MYSQL_DATABASE_USER):$(MYSQL_DATABASE_PASSWORD)@mariadb/$(MYSQL_DATABASE_NAME)
```

Your file should now look like this:

```
...
containers:
- image: quay.io/acend/example-web-python:latest
  imagePullPolicy: Always
  name: example-web-app
...
env:
- name: MYSQL_DATABASE_NAME
  valueFrom:
    secretKeyRef:
      key: database-name
      name: mariadb
- name: MYSQL_DATABASE_PASSWORD
  valueFrom:
    secretKeyRef:
      key: database-password
      name: mariadb
- name: MYSQL_DATABASE_ROOT_PASSWORD
  valueFrom:
    secretKeyRef:
      key: database-root-password
      name: mariadb
- name: MYSQL_DATABASE_USER
  valueFrom:
    secretKeyRef:
      key: database-user
      name: mariadb
- name: MYSQL_URI
  value: mysql://$(MYSQL_DATABASE_USER):$(MYSQL_DATABASE_PASSWORD)@mariadb/$(MYSQL_DATABASE_NAME)
```

Then use:

```
kubectl apply -f deployment_example-web-app.yaml --namespace <namespace>
```

to apply the changes.

The environment can also be checked with the `set env` command and the `--list` parameter:

- acend gmbh

```
kubectl set env deploy/example-web-app --list --namespace <namespace>
```

This will show the environment as follows:

```
# deployments/example-web-app, container example-web-app
# MYSQL_DATABASE_PASSWORD from secret mariadb, key database-password
# MYSQL_DATABASE_ROOT_PASSWORD from secret mariadb, key database-root-password
# MYSQL_DATABASE_USER from secret mariadb, key database-user
# MYSQL_DATABASE_NAME from secret mariadb, key database-name
MYSQL_URI=mysql://$(MYSQL_DATABASE_USER):$(MYSQL_DATABASE_PASSWORD)@mariadb/$(MYSQL_DATABASE_NAME)
```

Warning

Do not proceed with the lab before all example-web-app pods are restarted successfully.

The change of the deployment definition (environment change) triggers a new rollout and all example-web-app pods will be restarted. The application will not be connected to the database until all pods are restarted successfully.

In order to find out if the change worked we can either look at the container's logs (`kubectl logs <pod>`) or we could register some "Hellos" in the application, delete the Pod, wait for the new Pod to be started and check if they are still there.

Note

This does not work if we delete the database Pod as its data is not yet persisted.

Task 7.4: Manual database connection

As described in *6. Troubleshooting* we can log into a Pod with `kubectl exec -it <pod> -- /bin/bash`.

Show all Pods:

```
kubectl get pods --namespace <namespace>
```

Which gives you an output similar to this:

NAME	READY	STATUS	RESTARTS	AGE
example-web-app-574544fd68-qfkcm	1/1	Running	0	2m20s
mariadb-f845ccdb7-hf2x5	1/1	Running	0	31m
mariadb-1-deploy	0/1	Completed	0	11m

Log into the MariaDB Pod:

Note

As mentioned in *6. Troubleshooting*, remember to append the command with `winpty` if you're using Git Bash on Windows.

- acend gmbh

```
kubectl exec -it deployments/mariadb --namespace <namespace> -- /bin/bash
```

You are now able to connect to the database and display the data. Login with:

```
mysql -u$MYSQL_USER -p$MYSQL_PASSWORD -h$MARIADB_SERVICE_HOST $MYSQL_DATABASE
```

```
Welcome to the MariaDB monitor.  Commands end with ; or \g.
Your MariaDB connection id is 52810
Server version: 10.2.22-MariaDB MariaDB Server

Copyright (c) 2000, 2018, Oracle, MariaDB Corporation Ab and others.

Type 'help;' or '\h' for help. Type '\c' to clear the current input statement.

MariaDB [acend_exampledb]>
```

Show all tables with:

```
show tables;
```

Show any entered “Hellos” with:

```
select * from hello;
```

Task 7.5: Import a database dump

Our task is now to import this [dump.sql](#) into the MariaDB database running as a Pod. Use the `mysql` command line utility to do this. Make sure the database is empty beforehand. You could also delete and recreate the database.

Note

You can also copy local files into a Pod using `kubectl cp`. Be aware that the `tar` binary has to be present inside the container and on your operating system in order for this to work! Install `tar` on UNIX systems with e.g. your package manager, on Windows there's e.g. [cwRsync](#). If you cannot install `tar` on your host, there's also the possibility of logging into the Pod and using `curl -O <url>`.

Solution

This is how you copy the database dump into the MariaDB Pod.

Download the [dump.sql](#) or get it with `curl`:

```
curl -O https://raw.githubusercontent.com/acend/kubernetes-basics-training/main/content/en/docs/attaching-a-database/dump.sql
```

- acend gmbh

Copy the dump into the MariaDB Pod:

```
kubectl cp ./dump.sql <podname>:/tmp/ --namespace <namespace>
```

This is how you log into the MariaDB Pod:

```
kubectl exec -it <podname> --namespace <namespace> -- /bin/bash
```

This command shows how to drop the whole database:

```
mysql -u$MYSQL_USER -p$MYSQL_PASSWORD -h$MARIADB_SERVICE_HOST $MYSQL_DATABASE
```

```
drop database `acend_exampledb`;  
create database `acend_exampledb`;  
exit
```

Import a dump:

```
mysql -u$MYSQL_USER -p$MYSQL_PASSWORD -h$MARIADB_SERVICE_HOST $MYSQL_DATABASE < /tmp/dump.sql
```

Check your app to see the imported “Hello”.

Note

You can find your app URL by looking at your ingress:

```
kubectl get ingress --namespace <namespace>
```

Note

A database dump can be created as follows:

```
kubectl exec -it <podname> --namespace <namespace> -- /bin/bash
```

```
mysqldump --user=$MYSQL_USER --password=$MYSQL_PASSWORD -h$MARIADB_SERVICE_HOST $MYSQL_DATABASE > /tmp/dump.sql
```

```
kubectl cp <podname>:/tmp/dump.sql /tmp/dump.sql
```

8. Persistent storage

By default, data in containers is not persistent as was the case e.g. in *7. Attaching a database*. This means that the data written in a container is lost as soon as it does not exist anymore. We want to prevent this from happening. One possible solution to this problem is to use persistent storage.

Request storage

Attaching persistent storage to a Pod happens in two steps. The first step includes the creation of a so-called *PersistentVolumeClaim* (PVC) in our namespace. This claim defines amongst other things what size we would like to get.

The *PersistentVolumeClaim* only represents a request but not the storage itself. It is automatically going to be bound to a *PersistentVolume* by Kubernetes, one that has at least the requested size. If only volumes exist that have a bigger size than was requested, one of these volumes is going to be used. The claim will automatically be updated with the new size. If there are only smaller volumes available, the claim cannot be fulfilled as long as no volume with the exact same or larger size is created.

Attaching a volume to a Pod

In a second step, the PVC from before is going to be attached to the Pod. In *5. Scaling* we edited the deployment configuration in order to insert a readiness probe. We are now going to do the same for inserting the persistent volume.

Task 8.1: Add a PersistentVolume

The following command creates a *PersistentVolumeClaim* which requests a volume of 1Gi size. Save it to `pvc.yaml` :

```
apiVersion: v1
kind: PersistentVolumeClaim
metadata:
  name: mariadb-data
spec:
  accessModes:
    - ReadWriteOnce
  resources:
    requests:
      storage: 1Gi
```

And create it with:

```
kubectl apply -f pvc.yaml --namespace <namespace>
```

We now have to insert the volume definition in the correct section of the MariaDB deployment.

Change your local `mariadb.yaml` file and add the `volumeMounts` and `volumes` parts:

- acend gmbh

```
resources:
  limits:
    cpu: 500m
    memory: 512Mi
  requests:
    cpu: 50m
    memory: 128Mi
# start to copy here
volumeMounts:
- name: mariadb-data
  mountPath: /var/lib/mysql
volumes:
- name: mariadb-data
  persistentVolumeClaim:
    claimName: mariadb-data
```

Then apply the change with:

```
kubectl apply -f mariadb.yaml --namespace <namespace>
```

Note

Because we just changed the Deployment a new Pod was automatically redeployed. This unfortunately also means that we just lost the data we inserted before.

We need to redeploy the application pod, our application automatically creates the database schema at startup time. Wait for the database pod to be started fully before restarting the application pod.

If you want to force a redeployment of a Pod, you can use this:

```
kubectl rollout restart deployment example-web-app --namespace <namespace>
```

Using the command `kubectl get persistentvolumeclaim` or `kubectl get pvc`, we can display the freshly created PersistentVolumeClaim:

```
kubectl get pvc --namespace <namespace>
```

Which gives you an output similar to this:

NAME	STATUS	VOLUME	CAPACITY	ACCESS MODES	STORAGECLASS	AGE
mariadb-data	Bound	pvc-2cb78deb-d157-11e8-a406-42010a840034	1Gi	RWO	standard	11s

The two columns `STATUS` and `VOLUME` show us that our claim has been bound to the PersistentVolume `pvc-2cb78deb-d157-11e8-a406-42010a840034`.

Error case

If the container is not able to start it is the right moment to debug it! Check the logs from the container and search for the error.

- acend gmbh

```
kubectl logs mariadb-f845ccdb7-hf2x5 --namespace <namespace>
```

Note

If the container won't start because the data directory already has files in it, use the `kubectl exec` command mentioned in *7. Attaching a database* to check its content and remove it if necessary.

Task 8.2: Persistence check

Restore data

Repeat [the task to import a database dump](#) .

Test

Scale your MariaDB Pod to 0 replicas and back to 1. Observe that the new Pod didn't loose any data.

9. Additional concepts

Kubernetes does not only know Pods, Deployments, Services, etc. There are various other kinds of resources. In the next few labs, we are going to have a look at some of them.

9.1. StatefulSets

Stateless applications or applications with a stateful backend can be described as Deployments. However, sometimes your application has to be stateful. Examples would be an application that needs a static, non-changing hostname every time it starts or a clustered application with a strict start/stop order of its services (e.g. RabbitMQ). These features are offered by StatefulSets.

Note

This lab does not depend on other labs.

Consistent hostnames

While in normal Deployments a hash-based name of the Pods (also represented as the hostname inside the Pod) is generated, StatefulSets create Pods with preconfigured names. An example of a RabbitMQ cluster with three instances (Pods) could look like this:

```
rabbitmq-0  
rabbitmq-1  
rabbitmq-2
```

Scaling

Scaling is handled differently in StatefulSets. When scaling up from 3 to 5 replicas in a Deployment, two additional Pods are started at the same time (based on the configuration). Using a StatefulSet, scaling is done serially:

Let's use our RabbitMQ example again:

1. The StatefulSet is scaled up using: `kubectl scale deployment rabbitmq --replicas=5 --namespace <namespace>`
2. `rabbitmq-3` is started
3. As soon as Pod `rabbitmq-3` is in `Ready` state the same procedure starts for `rabbitmq-4`

When scaling down, the order is inverted. The highest-numbered Pod will be stopped first. As soon as it has finished terminating the now highest-numbered Pod is stopped. This procedure is repeated as long as the desired number of replicas has not been reached.

Update procedure

During an update of an application with a StatefulSet the highest-numbered Pod will be the first to be updated and only after a successful start the next Pod follows.

1. Highest-numbered Pod is stopped
2. New Pod (with new image tag) is started
3. If the new Pod successfully starts, the procedure is repeated for the second highest-numbered Pod
4. And so on

- acend gmbh

If the start of a new Pod fails, the update will be interrupted so that the architecture of your application won't break.

Dedicated persistent volumes

A very convenient feature is that unlike a Deployment a StatefulSet makes it possible to attach a different, dedicated persistent volume to each of its Pods. This is done using a so-called *VolumeClaimTemplate*. This spares you from defining identical Deployments with 1 replica each but different volumes.

Conclusion

The controllable and predictable behavior can be a perfect match for applications such as RabbitMQ or etcd, as you need unique names for such application clusters.

Task 9.1.1: Create a StatefulSet

Create a file named `sts_nginx-cluster.yaml` with the following definition of a StatefulSet:

```
apiVersion: apps/v1
kind: StatefulSet
metadata:
  name: nginx-cluster
spec:
  serviceName: "nginx"
  replicas: 1
  selector:
    matchLabels:
      app: nginx
  template:
    metadata:
      labels:
        app: nginx
    spec:
      containers:
        - name: nginx
          image: nginxinc/nginx-unprivileged:1.18-alpine
          ports:
            - containerPort: 8080
              name: nginx
          resources:
            limits:
              cpu: 40m
              memory: 64Mi
            requests:
              cpu: 10m
              memory: 32Mi
```

Create the StatefulSet:

```
kubectl apply -f sts_nginx-cluster.yaml --namespace <namespace>
```

To watch the pods' progress, open a second console and execute the watch command:

- acend gmbh

```
kubectl get pods --selector app=nginx -w --namespace <namespace>
```

Note

Friendly reminder that the `kubectl get -w` command will never end unless you terminate it with `CTRL-C`.

Task 9.1.2: Scale the StatefulSet

Scale the StatefulSet up:

```
kubectl scale statefulset nginx-cluster --replicas=3 --namespace <namespace>
```

You can again watch the pods' progress like you did in the first task.

Task 9.1.3: Update the StatefulSet

In order to update the image tag in use in a StatefulSet, you can use the `kubectl set image` command. Set the StatefulSet's image tag to `latest` :

```
kubectl set image statefulset nginx-cluster nginx=docker.io/nginxinc/nginx-unprivileged:latest --namespace <namespace>
```

Task 9.1.4: Rollback

Imagine you just realized that switching to the `latest` image tag was an awful idea (because it is generally not advisable). Rollback the change:

```
kubectl rollout undo statefulset nginx-cluster --namespace <namespace>
```

Task 9.1.5: Cleanup

As with every other Kubernetes resource you can delete the StatefulSet with:

Warning

To avoid issues on your personal progress dashboard, we would advise not to delete the StatefulSet from this lab

```
kubectl delete statefulset nginx-cluster --namespace <namespace>
```

Further information can be found in the [Kubernetes' StatefulSet documentation](#) or this [published article](#) .

9.2. DaemonSets

A DaemonSet is almost identical to a normal Deployment. The difference is that it makes sure that exactly one Pod is running on every (or some specified) Node. When a new Node is added, the DaemonSet automatically deploys a Pod on the new Node if its selector matches. When the DaemonSet is deleted, all related Pods are deleted.

One obvious use case for a DaemonSet is some kind of agent or daemon to e.g. grab logs from each Node of the cluster (e.g., Fluentd, Logstash or a Splunk forwarder).

More information about DaemonSet can be found in the [Kubernetes DaemonSet Documentation](#) .

9.3. CronJobs and Jobs

Jobs are different from normal Deployments: Jobs execute a time-constrained operation and report the result as soon as they are finished; think of a batch job. To achieve this, a Job creates a Pod and runs a defined command. A Job isn't limited to creating a single Pod, it can also create multiple Pods. When a Job is deleted, the Pods started (and stopped) by the Job are also deleted.

For example, a Job is used to ensure that a Pod is run until its completion. If a Pod fails, for example because of a Node error, the Job starts a new one. A Job can also be used to start multiple Pods in parallel.

More detailed information can be retrieved from the [Kubernetes documentation](#).

Note

This lab depends on *7. Attaching a database* or *8. Persistent storage*.

Task 9.3.1: Create a Job for a database dump

Similar to [the task to import a database dump](#), we now want to create a dump of the running database, but without the need of interactively logging into the Pod.

Let's first look at the Job resource that we want to create.

- acend gmbh

```
apiVersion: batch/v1
kind: Job
metadata:
  name: database-dump
spec:
  template:
    spec:
      containers:
      - name: mariadb
        image: mariadb:10.5
        command:
        - 'bash'
        - '-eo'
        - 'pipefail'
        - '-c'
        - >
          trap "echo Backup failed; exit 0" ERR;
          FILENAME=backup-`${MYSQL_DATABASE}`-`date +%Y-%m-%d_%H%M%S`.sql.gz;
          mysqldump --user=${MYSQL_USER} --password=${MYSQL_PASSWORD} --host=${MYSQL_HOST} --port=${MYSQL_PORT} --skip-
lock-tables --quick --add-drop-database --routines ${MYSQL_DATABASE} | gzip > /tmp/$FILENAME;
          echo "";
          echo "Backup successful"; du -h /tmp/$FILENAME;
        env:
        - name: MYSQL_DATABASE
          valueFrom:
            secretKeyRef:
              key: database-name
              name: mariadb
        - name: MYSQL_USER
          valueFrom:
            secretKeyRef:
              key: database-user
              name: mariadb
        - name: MYSQL_HOST
          value: mariadb
        - name: MYSQL_PORT
          value: "3306"
        - name: MYSQL_PASSWORD
          valueFrom:
            secretKeyRef:
              key: database-password
              name: mariadb
      resources:
        limits:
          cpu: 100m
          memory: 128Mi
        requests:
          cpu: 20m
          memory: 64Mi
      restartPolicy: Never
```

The parameter `.spec.template.spec.containers[0].image` shows that we use the same image as the running database. In contrast to the database Pod, we don't start a database afterwards, but run a `mysqldump` command, specified with `.spec.template.spec.containers[0].command`. To perform the dump, we use the environment variables of the database deployment to set the hostname, user and password parameters of the `mysqldump` command. The `MYSQL_PASSWORD` variable refers to the value of the secret, which is already used for the database Pod. This way we ensure that the dump is performed with the same credentials.

Let's create our Job: Create a file named `job_database-dump.yaml` with the content above and execute the following command:

```
kubectl apply -f ./job_database-dump.yaml --namespace <namespace>
```

Check if the Job was successful:

- acend gmbh

```
kubectl describe jobs/database-dump --namespace <namespace>
```

The executed Pod can be shown as follows:

```
kubectl get pods --namespace <namespace>
```

To show all Pods belonging to a Job in a human-readable format, the following command can be used:

```
kubectl get pods --selector=job-name=database-dump --output=go-template="{{range .items}}{{.metadata.name}}{{end}}" --namespace <namespace>
```

CronJobs

A CronJob is nothing else than a resource which creates a Job at a defined time, which in turn starts (as we saw in the previous section) a Pod to run a command. Typical use cases are cleanup Jobs, which tidy up old data for a running Pod, or a Job to regularly create and save a database dump as we just did during this lab.

The CronJob's definition will remind you of the Deployment's structure, or really any other control resource. There's most importantly the `schedule` specification in [cron schedule format](#), some more things you could define and then the Job's definition itself that is going to be created by the CronJob:

```
apiVersion: batch/v1
kind: CronJob
metadata:
  name: database-dump
spec:
  schedule: "5 4 * * *"
  concurrencyPolicy: "Replace"
  startingDeadlineSeconds: 200
  successfulJobsHistoryLimit: 3
  failedJobsHistoryLimit: 1
  jobTemplate:
    spec:
      template:
        spec:
          containers:
            - name: mariadb
              ...
```

Further information can be found in the [Kubernetes CronJob documentation](#).

9.4. ConfigMaps

Similar to environment variables, *ConfigMaps* allow you to separate the configuration for an application from the image. Pods can access those variables at runtime which allows maximum portability for applications running in containers. In this lab, you will learn how to create and use ConfigMaps.

ConfigMap creation

A ConfigMap can be created using the `kubectl create configmap` command as follows:

```
kubectl create configmap <name> <data-source> --namespace <namespace>
```

Where the `<data-source>` can be a file, directory, or command line input.

Task 9.4.1: Java properties as ConfigMap

A classic example for ConfigMaps are properties files of Java applications which can't be configured with environment variables.

First, create a file called `java.properties` with the following content:

```
key=value  
key2=value2
```

Now you can create a ConfigMap based on that file:

```
kubectl create configmap javaconfiguration --from-file=./java.properties --namespace <namespace>
```

Verify that the ConfigMap was created successfully:

```
kubectl get configmaps --namespace <namespace>
```

NAME	DATA	AGE
javaconfiguration	1	7s

Have a look at its content:

```
kubectl get configmap javaconfiguration -o yaml --namespace <namespace>
```

Which should yield output similar to this one:

- acend gmbh

```
apiVersion: v1
kind: ConfigMap
metadata:
  name: javaconfiguration
data:
  java.properties: |
    key=value
    key2=value2
```

Task 9.4.2: Attach the ConfigMap to a container

Next, we want to make a ConfigMap accessible for a container. There are basically the following possibilities to achieve [this](#) :

- ConfigMap properties as environment variables in a Deployment
- Command line arguments via environment variables
- Mounted as volumes in the container

In this example, we want the file to be mounted as a volume inside the container.

Basically, a Deployment has to be extended with the following config:

```
...
volumeMounts:
- mountPath: /etc/config
  name: config-volume
...
volumes:
- configMap:
  defaultMode: 420
  name: javaconfiguration
  name: config-volume
...
```

Here is a complete example Deployment of a sample Java app:

- acend gmbh

```
apiVersion: apps/v1
kind: Deployment
metadata:
  labels:
    app: spring-boot-example
    name: spring-boot-example
spec:
  progressDeadlineSeconds: 600
  replicas: 1
  revisionHistoryLimit: 10
  selector:
    matchLabels:
      app: spring-boot-example
  strategy:
    rollingUpdate:
      maxSurge: 25%
      maxUnavailable: 25%
    type: RollingUpdate
  template:
    metadata:
      labels:
        app: spring-boot-example
    spec:
      containers:
        - image: appuio/example-spring-boot
          imagePullPolicy: Always
          name: example-spring-boot
          resources:
            limits:
              cpu: 1
              memory: 768Mi
            requests:
              cpu: 20m
              memory: 32Mi
          terminationMessagePath: /dev/termination-log
          terminationMessagePolicy: File
          volumeMounts:
            - mountPath: /etc/config
              name: config-volume
      dnsPolicy: ClusterFirst
      restartPolicy: Always
      schedulerName: default-scheduler
      securityContext: {}
      terminationGracePeriodSeconds: 30
      volumes:
        - configMap:
            defaultMode: 420
            name: javaconfiguration
          name: config-volume
```

This means that the container should now be able to access the ConfigMap's content in /etc/config/java.properties . Let's check:

```
kubectl exec <pod> --namespace <namespace> -- cat /etc/config/java.properties
```

Note

On Windows, you can use Git Bash with `winpty kubectl exec -it <pod> --namespace <namespace> -- cat //etc/config/java.properties.`

```
key=value
key2=value2
```

- acend gmbh

Like this, the property file can be read and used by the application inside the container. The image stays portable to other environments.

Task 9.4.3: ConfigMap environment variables

Use a ConfigMap by [populating environment variables into the container](#) instead of a file.

9.5. ResourceQuotas and LimitRanges

In this lab, we are going to look at ResourceQuotas and LimitRanges. As Kubernetes users, we are most certainly going to encounter the limiting effects that ResourceQuotas and LimitRanges impose.

Warning

For this lab to work it is vital that you use the namespace `<username>-quota` !

ResourceQuotas

ResourceQuotas among other things limit the amount of resources Pods can use in a Namespace. They can also be used to limit the total number of a certain resource type in a Namespace. In more detail, there are these kinds of quotas:

- *Compute ResourceQuotas* can be used to limit the amount of memory and CPU
- *Storage ResourceQuotas* can be used to limit the total amount of storage and the number of PersistentVolumeClaims, generally or specific to a StorageClass
- *Object count quotas* can be used to limit the number of a certain resource type such as Services, Pods or Secrets

Defining ResourceQuotas makes sense when the cluster administrators want to have better control over consumed resources. A typical use case are public offerings where users pay for a certain guaranteed amount of resources which must not be exceeded.

In order to check for defined quotas in your Namespace, simply see if there are any of type ResourceQuota:

```
kubectl get resourcequota --namespace <namespace>-quota
```

To show in detail what kinds of limits the quota imposes:

```
kubectl describe resourcequota <quota-name> --namespace <namespace>-quota
```

For more details, have look at [Kubernetes' documentation about resource quotas](#) .

Requests and limits

As we've already seen, compute ResourceQuotas limit the amount of memory and CPU we can use in a Namespace. Only defining a ResourceQuota, however is not going to have an effect on Pods that don't define the amount of resources they want to use. This is where the concept of limits and requests comes into play.

Limits and requests on a Pod, or rather on a container in a Pod, define how much memory and CPU this container wants to consume at least (request) and at most (limit). Requests mean that the container will be guaranteed to get at least this amount of resources, limits represent the upper boundary which cannot be crossed. Defining these values helps Kubernetes in determining on which Node to schedule the Pod because it knows how many resources should be available for it.

Note

Containers using more CPU time than what their limit allows will be throttled. Containers using more memory than what they are allowed to use will be killed.

Defining limits and requests on a Pod that has one container looks like this:

```
apiVersion: v1
kind: Pod
metadata:
  name: lr-demo
  namespace: lr-example
spec:
  containers:
  - name: lr-demo-ctr
    image: docker.io/nginxinc/nginx-unprivileged:latest
    resources:
      limits:
        memory: "200Mi"
        cpu: "700m"
      requests:
        memory: "200Mi"
        cpu: "700m"
```

You can see the familiar binary unit “Mi” is used for the memory value. Other binary (“Gi”, “Ki”, ...) or decimal units (“M”, “G”, “K”, ...) can be used as well.

The CPU value is denoted as “m”. “m” stands for *millicpu* or sometimes also referred to as *millicores* where “1000m” is equal to one core/vCPU/hyperthread.

Quality of service

Setting limits and requests on containers has yet another effect: It might change the Pod’s *Quality of Service* class. There are three such *QoS* classes:

- *Guaranteed*
- *Burstable*
- *BestEffort*

The Guaranteed QoS class is applied to Pods that define both limits and requests for both memory and CPU resources on all their containers. The most important part is that each request has the same value as the limit. Pods that belong to this QoS class will never be killed by the scheduler because of resources running out on a Node.

Note

If a container only defines its limits, Kubernetes automatically assigns a request that matches the limit.

The Burstable QoS class means that limits and requests on a container are set, but they are different. It is enough to define limits and requests on one container of a Pod even though there might be more, and it also only has to define limits and requests on memory or CPU, not necessarily both.

The BestEffort QoS class applies to Pods that do not define any limits and requests at all on any containers. As its class name suggests, these are the kinds of Pods that will be killed by the scheduler first if a Node runs out of memory or CPU. As you might have already guessed by now, if there are no BestEffort QoS Pods, the scheduler will begin to kill Pods belonging to the class of *Burstable*. A Node hosting only Pods of class Guaranteed will (theoretically) never run out of resources.

For more examples have a look at the [Kubernetes documentation about Quality of Service](#) .

LimitRanges

As you now know what limits and requests are, we can come back to the statement made above:

As we've already seen, compute ResourceQuotas limit the amount of memory and CPU we can use in a Namespace. Only defining a ResourceQuota, however is not going to have an effect on Pods that don't define the amount of resources they want to use. This is where the concept of limits and requests comes into play.

So, if a cluster administrator wanted to make sure that every Pod in the cluster counted against the compute ResourceQuota, the administrator would have to have a way of defining some kind of default limits and requests that were applied if none were defined in the containers. This is exactly what *LimitRanges* are for.

Quoting the [Kubernetes documentation](#) , LimitRanges can be used to:

- Enforce minimum and maximum compute resource usage per Pod or container in a Namespace
- Enforce minimum and maximum storage requests per PersistentVolumeClaim in a Namespace
- Enforce a ratio between request and limit for a resource in a Namespace
- Set default request/limit for compute resources in a Namespace and automatically inject them to containers at runtime

If for example a container did not define any requests or limits and there was a LimitRange defining the default values, these default values would be used when deploying said container. However, as soon as limits or requests were defined, the default values would no longer be applied.

The possibility of enforcing minimum and maximum resources and defining ResourceQuotas per Namespace allows for many combinations of resource control.

Task 9.5.1: Namespace

Warning

Remember to use the namespace `<username>-quota` , otherwise this lab will not work!

Analyse the LimitRange in your Namespace (there has to be one, if not you are using the wrong Namespace):

```
kubectl describe limitrange --namespace <namespace>-quota
```

The command above should output this (name and Namespace will vary):

Name:	ce01a1b6-a162-479d-847c-4821255cc6db						
Namespace:	eltony-quota-lab						
Type	Resource	Min	Max	Default Request	Default Limit	Max Limit/Request Ratio	
-----	-----	---	---	-----	-----	-----	
Container	memory	-	-	16Mi	32Mi	-	
Container	cpu	-	-	10m	100m	-	

Check for the ResourceQuota in your Namespace (there has to be one, if not you are using the wrong Namespace):

- acend gmbh

```
kubectl describe quota --namespace <namespace>-quota
```

The command above will produce an output similar to the following (name and namespace may vary)

Name:	lab-quota	
Namespace:	eltony-quota-lab	
Resource	Used	Hard
-----	----	----
requests.cpu	0	100m
requests.memory	0	100Mi

Task 9.5.2: Default memory limit

Create a Pod using the stress image:

```
apiVersion: v1
kind: Pod
metadata:
  name: stress2much
spec:
  containers:
    - command:
      - stress
      - --vm
      - "1"
      - --vm-bytes
      - 85M
      - --vm-hang
      - "1"
      image: quay.io/acend/stress:latest
      imagePullPolicy: Always
      name: stress
```

Apply this resource with:

```
kubectl apply -f pod_stress2much.yaml --namespace <namespace>-quota
```

Note

You have to actively terminate the following command pressing **CTRL+C** on your keyboard.

Watch the Pod's creation with:

```
kubectl get pods --watch --namespace <namespace>-quota
```

You should see something like the following:

- acend gmbh

NAME	READY	STATUS	RESTARTS	AGE
stress2much	0/1	ContainerCreating	0	1s
stress2much	0/1	ContainerCreating	0	2s
stress2much	0/1	OOMKilled	0	5s
stress2much	1/1	Running	1	7s
stress2much	0/1	OOMKilled	1	9s
stress2much	0/1	CrashLoopBackOff	1	20s

The `stress2much` Pod was OOM (out of memory) killed. We can see this in the `STATUS` field. Another way to find out why a Pod was killed is by checking its status. Output the Pod's YAML definition:

```
kubectl get pod stress2much --output yaml --namespace <namespace>-quota
```

Near the end of the output you can find the relevant status part:

```
containerStatuses:
- containerID: docker://da2473f1c8ccdfbb824d03689e9fe738ed689853e9c2643c37f206d10f93a73
  image: quay.io/acend/stress:latest
  lastState:
    terminated:
      ...
      reason: OOMKilled
      ...
```

So let's look at the numbers to verify the container really had too little memory. We started the `stress` command using the parameter `--vm-bytes 85M` which means the process wants to allocate 85 megabytes of memory. Again looking at the Pod's YAML definition with:

```
kubectl get pod stress2much --output yaml --namespace <namespace>-quota
```

reveals the following values:

```
...
  resources:
    limits:
      cpu: 100m
      memory: 32Mi
    requests:
      cpu: 10m
      memory: 16Mi
  ...
```

These are the values from the LimitRange, and the defined limit of 32 MiB of memory prevents the `stress` process of ever allocating the desired 85 MB.

Let's fix this by recreating the Pod and explicitly setting the memory request to 85 MB.

First, delete the `stress2much` pod with:

```
kubectl delete pod stress2much --namespace <namespace>-quota
```

- acend gmbh

Then create a new Pod where the requests and limits are set:

```
apiVersion: v1
kind: Pod
metadata:
  name: stress
spec:
  containers:
    - command:
      - stress
      - --vm
      - "1"
      - --vm-bytes
      - 85M
      - --vm-hang
      - "1"
      image: quay.io/acend/stress:latest
      imagePullPolicy: Always
      name: stress
    resources:
      limits:
        cpu: 100m
        memory: 100Mi
      requests:
        cpu: 10m
        memory: 85Mi
```

And apply this again with:

```
kubectl apply -f pod_stress.yaml --namespace <namespace>-quota
```

Note

Remember, if you only set the limit, the request will be set to the same value.

You should now see that the Pod is successfully running:

NAME	READY	STATUS	RESTARTS	AGE
stress	1/1	Running	0	25s

Task 9.5.3: Hitting the quota

Create another Pod, again using the `stress` image. This time our application is less demanding and only needs 10 MB of memory (`--vm-bytes 10M`):

Create a new Pod resource with:

- acend gmbh

```
apiVersion: v1
kind: Pod
metadata:
  name: overbooked
spec:
  containers:
    - command:
      - stress
      - --vm
      - "1"
      - --vm-bytes
      - 10M
      - --vm-hang
      - "1"
      image: quay.io/acend/stress:latest
      imagePullPolicy: Always
      name: overbooked
```

```
kubectl apply -f pod_overbooked.yaml --namespace <namespace>-quota
```

We are immediately confronted with an error message:

```
Error from server (Forbidden): pods "overbooked" is forbidden: exceeded quota: lab-quota, requested: memory=16Mi, used:
memory=85Mi, limited: memory=100Mi
```

The default request value of 16 MiB of memory that was automatically set on the Pod lets us hit the quota which in turn prevents us from creating the Pod.

Let's have a closer look at the quota with:

```
kubectl get quota --output yaml --namespace <namespace>-quota
```

which should output the following YAML definition:

```
...
status:
  hard:
    cpu: 100m
    memory: 100Mi
  used:
    cpu: 20m
    memory: 80Mi
...
```

The most interesting part is the quota's status which reveals that we cannot use more than 100 MiB of memory and that 80 MiB are already used.

Fortunately, our application can live with less memory than what the LimitRange sets. Let's set the request to the remaining 10 MiB:

- acend gmbh

```
apiVersion: v1
kind: Pod
metadata:
  name: overbooked
spec:
  containers:
    - command:
        - stress
        - --vm
        - "1"
        - --vm-bytes
        - 10M
        - --vm-hang
        - "1"
      image: quay.io/acend/stress:latest
      imagePullPolicy: Always
      name: overbooked
      resources:
        limits:
          cpu: 100m
          memory: 50Mi
        requests:
          cpu: 10m
          memory: 10Mi
```

And apply with:

```
kubectl apply -f pod_overbooked.yaml --namespace <namespace>-quota
```

Even though the limits of both Pods combined overstretch the quota, the requests do not and so the Pods are allowed to run.

9.6. Init containers

A Pod can have multiple containers running apps within it, but it can also have one or more *init containers*, which are run before the app container is started.

Init containers are exactly like regular containers, except:

- Init containers always run to completion.
- Each init container must complete successfully before the next one starts.

Check [Init Containers](#) from the Kubernetes documentation for more details.

Task 9.6.1: Add an init container

In *7. Attaching a database* you created the `example-web-app` application. In this task, you are going to add an init container which checks if the MariaDB database is ready to be used before actually starting your example application.

Edit your existing `example-web-app` Deployment by changing your local `deployment_example-web-app.yaml`. Add the init container into the existing Deployment (same indentation level as containers):

```
...
spec:
  initContainers:
  - name: wait-for-db
    image: docker.io/busybox:1.28
    command:
    [
      "sh",
      "-c",
      "until nslookup mariadb.${cat /var/run/secrets/kubernetes.io/serviceaccount/namespace}.svc.cluster.local;
do echo waiting for mydb; sleep 2; done",
    ]
...

```

And then apply again with:

```
kubectl apply -f deployment_example-web-app.yaml --namespace <namespace>
```

Note

This obviously only checks if there is a DNS Record for your MariaDB Service and not if the database is ready. But you get the idea, right?

Let's see what has changed by analyzing your newly created `example-web-app` Pod with the following command (use `kubectl get pod` or auto-completion to get the Pod name):

```
kubectl describe pod <pod> --namespace <namespace>
```

You see the new init container with the name `wait-for-db` :

- acend gmbh

```
...
Init Containers:
  wait-for-db:
    Container ID:   docker://77e6e309c88cfe62d03ed97e8fae20704bbf547a1e717a8f699ba79d9879cca2
    Image:          busybox
    Image ID:       docker-pullable://busybox@sha256:141c253bc4c3fd0a201d32dc1f493bcf3fff003b6df416dea4f41046e0f37d47
    Port:           <none>
    Host Port:      <none>
    Command:
      sh
      -c
      until nslookup mariadb.$(cat /var/run/secrets/kubernetes.io/serviceaccount/namespace).svc.cluster.local; do echo
waiting for mydb; sleep 2; done
    State:          Terminated
      Reason:       Completed
      Exit Code:    0
      Started:      Tue, 10 Nov 2020 21:00:24 +0100
      Finished:     Tue, 10 Nov 2020 21:02:52 +0100
    Ready:          True
    Restart Count:  0
    Environment:    <none>
    Mounts:
      /var/run/secrets/kubernetes.io/serviceaccount from default-token-xz2b7 (ro)
...
```

The init container has the `State: Terminated` and an `Exit Code: 0` which means it was successful. That's what we wanted, the init container was successfully executed before our main application.

You can also check the logs of the init container with:

```
kubectl logs -c wait-for-db <pod> --namespace <namespace>
```

Which should give you something similar to:

```
Server:      10.43.0.10
Address 1:   10.43.0.10 kube-dns.kube-system.svc.cluster.local

Name:        mariadb.acend-test.svc.cluster.local
Address 1:   10.43.243.105 mariadb.acend-test.svc.cluster.local
```

Check [Init Container](#) from the Kubernetes documentation for more details.

9.7. Sidecar containers

Let's first have another look at the Pod's description [on the Kubernetes documentation page](#) :

A Pod (as in a pod of whales or pea pod) is a group of one or more containers (such as Docker containers), with shared storage/network, and a specification for how to run the containers. A Pod's contents are always co-located and co-scheduled, and run in a shared context. A Pod models an application-specific "logical host" - it contains one or more application containers which are relatively tightly coupled — in a pre-container world, being executed on the same physical or virtual machine would mean being executed on the same logical host. The shared context of a Pod is a set of Linux namespaces, cgroups, and potentially other facets of isolation - the same things that isolate a Docker container. Within a Pod's context, the individual applications may have further sub-isolations applied.

A sidecar container is a utility container in the Pod. Its purpose is to support the main container. It is important to note that the standalone sidecar container does not serve any purpose, it must be paired with one or more main containers. Generally, sidecar containers are reusable and can be paired with numerous types of main containers.

In a sidecar pattern, the functionality of the main container is extended or enhanced by a sidecar container without strong coupling between the two. Although it is always possible to build sidecar container functionality into the main container, there are several benefits with this pattern:

- Different resource profiles, i.e. independent resource accounting and allocation
- Clear separation of concerns at packaging level, i.e. no strong coupling between containers
- Reusability, i.e., sidecar containers can be paired with numerous "main" containers
- Failure containment boundary, making it possible for the overall system to degrade gracefully
- Independent testing, packaging, upgrade, deployment and if necessary rollback

Task 9.7.1: Add a Prometheus MySQL exporter as a sidecar

In *8. Persistent storage* you created a MariaDB deployment. In this task you are going to add the [Prometheus MySQL exporter](#) to it.

Change the existing `mariadb` Deployment by first editing your local `mariadb.yaml` file. Add a new (sidecar) container into your Deployment:

```
containers:
- ...
- image: docker.io/prom/mysqld-exporter:v0.14.0
  name: mysqld-exporter
  env:
  - name: MYSQL_DATABASE_ROOT_PASSWORD
    valueFrom:
      secretKeyRef:
        key: database-root-password
        name: mariadb
  - name: DATA_SOURCE_NAME
    value: root:${MYSQL_DATABASE_ROOT_PASSWORD}@localhost:3306)/
...

```

and then apply the change with:

```
kubectl apply -f mariadb.yaml --namespace <namespace>
```

- acend gmbh

Your Pod now has two running containers. Verify this with:

```
kubectl get pod --namespace <namespace>
```

The output should look similar to this:

NAME	READY	STATUS	RESTARTS	AGE
mariadb-65559644c9-cdjkk	2/2	Running	0	5m35s

Note the `READY` column which shows you 2 ready containers.

You can get the logs from the `mysqld-exporter` with:

```
kubectl logs <pod> -c mysqld-exporter --namespace <namespace>
```

Which gives you an output similar to this:

```
time="2020-05-10T11:31:02Z" level=info msg="Starting mysqld_exporter (version=0.12.1, branch=HEAD, revision=48667bf7c3b438b5e93b259f3d17b70a7c9aff96)" source="mysqld_exporter.go:257"
time="2020-05-10T11:31:02Z" level=info msg="Build context (go=go1.12.7, user=root@0b3e56a7bc0a, date=20190729-12:35:58)" source="mysqld_exporter.go:258"
time="2020-05-10T11:31:02Z" level=info msg="Enabled scrapers:" source="mysqld_exporter.go:269"
time="2020-05-10T11:31:02Z" level=info msg="--collect.global_variables" source="mysqld_exporter.go:273"
time="2020-05-10T11:31:02Z" level=info msg="--collect.slave_status" source="mysqld_exporter.go:273"
time="2020-05-10T11:31:02Z" level=info msg="--collect.global_status" source="mysqld_exporter.go:273"
time="2020-05-10T11:31:02Z" level=info msg="--collect.info_schema.query_response_time" source="mysqld_exporter.go:273"
time="2020-05-10T11:31:02Z" level=info msg="--collect.info_schema.innodb_cmp" source="mysqld_exporter.go:273"
time="2020-05-10T11:31:02Z" level=info msg="--collect.info_schema.innodb_cmpmem" source="mysqld_exporter.go:273"
time="2020-05-10T11:31:02Z" level=info msg="Listening on :9104" source="mysqld_exporter.go:283"
```

By using the `port-forward` subcommand, you can even have a look at the Prometheus metrics:

```
kubectl port-forward <pod> 9104 --namespace <namespace>
```

And then use `curl` to check the `mysqld_exporter` metrics with:

```
curl http://localhost:9104/metrics
```

10. Security

10.1. Network policies

Network Policies

One CNI function is the ability to enforce network policies and implement an in-cluster zero-trust container strategy. Network policies are a default Kubernetes object for controlling network traffic, but a CNI such as [Cilium](#) or [Calico](#) is required to enforce them. We will demonstrate traffic blocking with our simple app.

Note

If you are not yet familiar with Kubernetes Network Policies we suggest going to the [Kubernetes Documentation](#).

Task 10.1.1: Deploy a simple frontend/backend application

First we need a simple application to show the effects on Kubernetes network policies. Let's have a look at the following resource definitions:

```
---
apiVersion: apps/v1
kind: Deployment
metadata:
  name: frontend
  labels:
    app: frontend
spec:
  replicas: 1
  selector:
    matchLabels:
      app: frontend
  template:
    metadata:
      labels:
        app: frontend
    spec:
      containers:
        - name: frontend-container
          image: docker.io/byrnedo/alpine-curl:0.1.8
          imagePullPolicy: IfNotPresent
          command: [ "/bin/ash", "-c", "sleep 1000000000" ]
---
apiVersion: apps/v1
kind: Deployment
metadata:
  name: not-frontend
  labels:
    app: not-frontend
spec:
  replicas: 1
  selector:
    matchLabels:
      app: not-frontend
  template:
    metadata:
      labels:
        app: not-frontend
    spec:
      containers:
        - name: not-frontend-container
          image: docker.io/byrnedo/alpine-curl:0.1.8
```

```

        image: docker.io/tyrnedo/alpine-curl:0.1.8 - acend gmbh
        imagePullPolicy: IfNotPresent
        command: [ "/bin/ash", "-c", "sleep 1000000000" ]
---
apiVersion: apps/v1
kind: Deployment
metadata:
  name: backend
  labels:
    app: backend
spec:
  replicas: 1
  selector:
    matchLabels:
      app: backend
  template:
    metadata:
      labels:
        app: backend
    spec:
      containers:
        - name: backend-container
          env:
            - name: PORT
              value: "8080"
          ports:
            - containerPort: 8080
          image: docker.io/cilium/json-mock:1.2
          imagePullPolicy: IfNotPresent
---
apiVersion: v1
kind: Service
metadata:
  name: backend
  labels:
    app: backend
spec:
  type: ClusterIP
  selector:
    app: backend
  ports:
    - name: http
      port: 8080

```

The application consists of two client deployments (`frontend` and `not-frontend`) and one backend deployment (`backend`). We are going to send requests from the frontend and not-frontend pods to the backend pod.

Create a file `simple-app.yaml` with the above content.

Deploy the app:

```
kubectl apply -f simple-app.yaml
```

this gives you the following output:

```

deployment.apps/frontend created
deployment.apps/not-frontend created
deployment.apps/backend created
service/backend created

```

Verify with the following command that everything is up and running:

- acend gmbh

```
kubectl get all
```

NAME	READY	STATUS	RESTARTS	AGE
pod/backend-65f7c794cc-b9j66	1/1	Running	0	3m17s
pod/frontend-76fbb99468-mbzcw	1/1	Running	0	3m17s
pod/not-frontent-8f467ccbd-cbks8	1/1	Running	0	3m17s

NAME	TYPE	CLUSTER-IP	EXTERNAL-IP	PORT(S)	AGE
service/backend	ClusterIP	10.97.228.29	<none>	8080/TCP	3m17s
service/kubernetes	ClusterIP	10.96.0.1	<none>	443/TCP	45m

NAME	READY	UP-TO-DATE	AVAILABLE	AGE
deployment.apps/backend	1/1	1	1	3m17s
deployment.apps/frontend	1/1	1	1	3m17s
deployment.apps/not-frontent	1/1	1	1	3m17s

NAME	DESIRED	CURRENT	READY	AGE
replicaset.apps/backend-65f7c794cc	1	1	1	3m17s
replicaset.apps/frontend-76fbb99468	1	1	1	3m17s
replicaset.apps/not-frontent-8f467ccbd	1	1	1	3m17s

Let us make life a bit easier by storing the pods name into an environment variable so we can reuse it later again:

```
export FRONTEND=$(kubectl get pods -l app=frontend -o jsonpath='{.items[0].metadata.name}')
echo ${FRONTEND}
export NOT_FRONTEND=$(kubectl get pods -l app=not-frontent -o jsonpath='{.items[0].metadata.name}')
echo ${NOT_FRONTEND}
```

Task 10.1.2: Verify connectivity

Now we generate some traffic as a baseline test.

```
kubectl exec -ti ${FRONTEND} -- curl -I --connect-timeout 5 backend:8080
```

and

```
kubectl exec -ti ${NOT_FRONTEND} -- curl -I --connect-timeout 5 backend:8080
```

This will execute a simple `curl` call from the `frontend` and `not-frontent` application to the `backend` application:

- acend gmbh

```
# Frontend
HTTP/1.1 200 OK
X-Powered-By: Express
Vary: Origin, Accept-Encoding
Access-Control-Allow-Credentials: true
Accept-Ranges: bytes
Cache-Control: public, max-age=0
Last-Modified: Sat, 26 Oct 1985 08:15:00 GMT
ETag: W/"83d-7438674ba0"
Content-Type: text/html; charset=UTF-8
Content-Length: 2109
Date: Tue, 23 Nov 2021 12:50:44 GMT
Connection: keep-alive
```

```
# Not Frontend
HTTP/1.1 200 OK
X-Powered-By: Express
Vary: Origin, Accept-Encoding
Access-Control-Allow-Credentials: true
Accept-Ranges: bytes
Cache-Control: public, max-age=0
Last-Modified: Sat, 26 Oct 1985 08:15:00 GMT
ETag: W/"83d-7438674ba0"
Content-Type: text/html; charset=UTF-8
Content-Length: 2109
Date: Tue, 23 Nov 2021 12:50:44 GMT
Connection: keep-alive
```

and we see, both applications can connect to the `backend` application.

Until now ingress and egress policy enforcement are still disabled on all of our pods because no network policy has been imported yet selecting any of the pods. Let us change this.

Task 10.1.3: Deny traffic with a Network Policy

We block traffic by applying a network policy. Create a file `backend-ingress-deny.yaml` with the following content:

```
---
kind: NetworkPolicy
apiVersion: networking.k8s.io/v1
metadata:
  name: backend-ingress-deny
spec:
  podSelector:
    matchLabels:
      app: backend
  policyTypes:
    - Ingress
```

The policy will deny all ingress traffic as it is of type `Ingress` but specifies no allow rule, and will be applied to all pods with the `app=backend` label thanks to the `podSelector`.

Ok, then let's create the policy with:

```
kubectl apply -f backend-ingress-deny.yaml
```

and you can verify the created `NetworkPolicy` with:

- acend gmbh

```
kubectl get netpol
```

which gives you an output similar to this:

NAME	POD-SELECTOR	AGE
backend-ingress-deny	app=backend	2s

Task 10.1.4: Verify connectivity again

We can now execute the connectivity check again:

```
kubectl exec -ti ${FRONTEND} -- curl -I --connect-timeout 5 backend:8080
```

and

```
kubectl exec -ti ${NOT_FRONTEND} -- curl -I --connect-timeout 5 backend:8080
```

but this time you see that the frontend and not-frontend application cannot connect anymore to the backend :

```
# Frontend
curl: (28) Connection timed out after 5001 milliseconds
command terminated with exit code 28
# Not Frontend
curl: (28) Connection timed out after 5001 milliseconds
command terminated with exit code 28
```

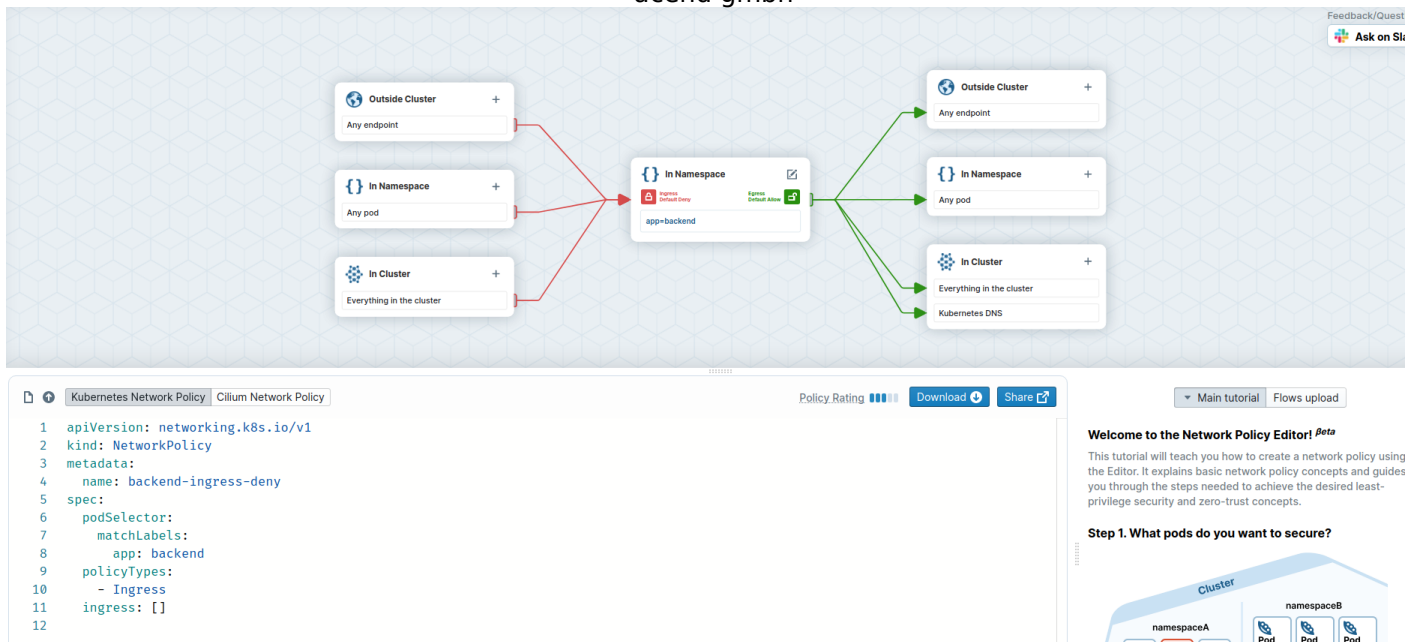
The network policy correctly switched the default ingress behavior from default allow to default deny.

Let's now selectively re-allow traffic again, but only from frontend to backend.

Task 10.1.5: Allow traffic from frontend to backend

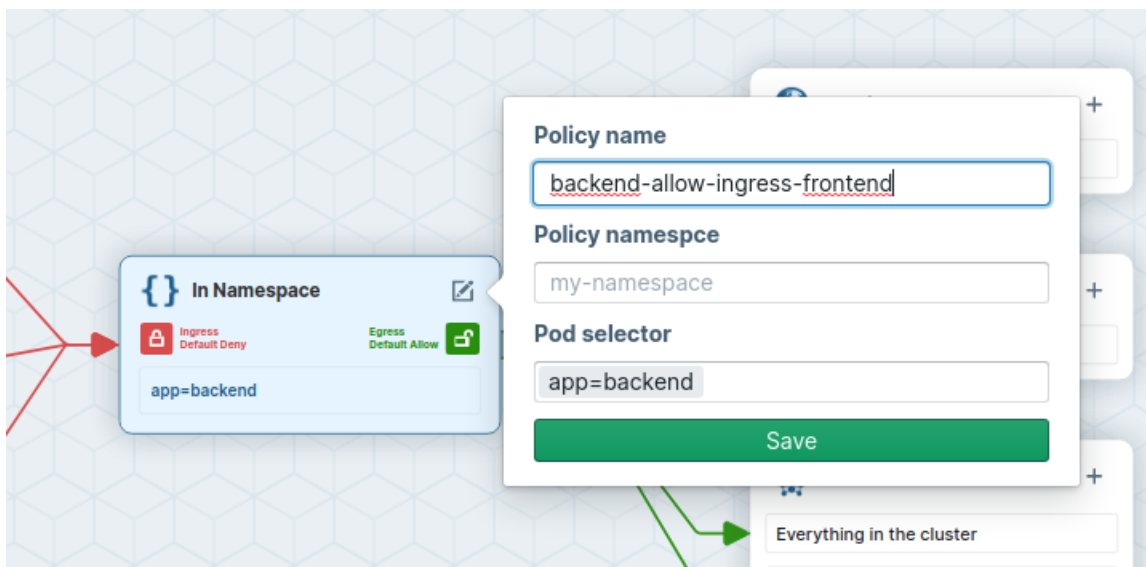
We can do it by crafting a new network policy manually, but we can also use the Network Policy Editor made by Cilium to help us out:

- acend gmbh



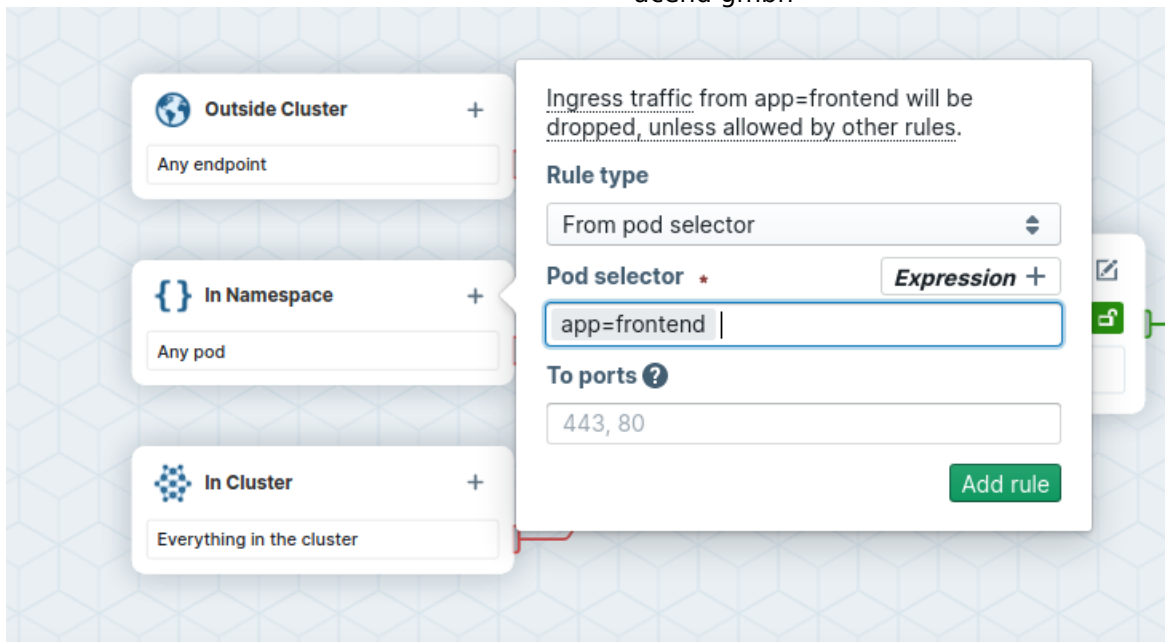
Above you see our original policy, we create an new one with the editor now.

- Go to <https://editor.cilium.io/>
- Name the network policy to backend-allow-ingress-frontend (using the Edit button in the center).
- add `app=backend` as Pod Selector
- Set Ingress to default deny

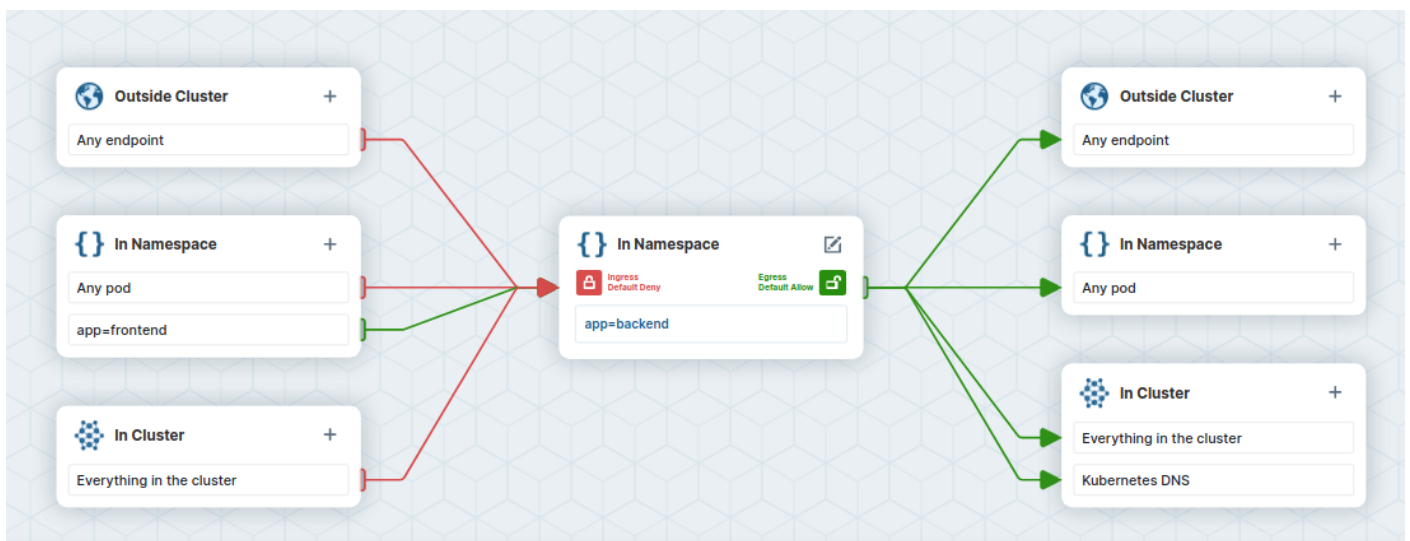


- On the ingress side, add `app=frontend` as podSelector for pods in the same Namespace.

- acend gmbh



- Inspect the ingress flow colors: the policy will deny all ingress traffic to pods labeled `app=backend`, except for traffic coming from pods labeled `app=frontend`.



- Copy the policy YAML into a file named `backend-allow-ingress-frontend.yaml`. Make sure to use the `NetworkPolicy` and not the `CiliumNetworkPolicy`!

The file should look like this:

- acend gmbh

```
---
apiVersion: networking.k8s.io/v1
kind: NetworkPolicy
metadata:
  name: "backend-allow-ingress-frontend"
spec:
  podSelector:
    matchLabels:
      app: backend
  policyTypes:
    - Ingress
  ingress:
    - from:
      - podSelector:
          matchLabels:
            app: frontend
```

Apply the new policy:

```
kubectl apply -f backend-allow-ingress-frontend.yaml
```

and then execute the connectivity test again:

```
kubectl exec -ti ${FRONTEND} -- curl -I --connect-timeout 5 backend:8080
```

and

```
kubectl exec -ti ${NOT_FRONTEND} -- curl -I --connect-timeout 5 backend:8080
```

This time, the `frontend` application is able to connect to the `backend` but the `not-frontend` application still cannot connect to the `backend` :

```
# Frontend
HTTP/1.1 200 OK
X-Powered-By: Express
Vary: Origin, Accept-Encoding
Access-Control-Allow-Credentials: true
Accept-Ranges: bytes
Cache-Control: public, max-age=0
Last-Modified: Sat, 26 Oct 1985 08:15:00 GMT
ETag: W/"83d-7438674ba0"
Content-Type: text/html; charset=UTF-8
Content-Length: 2109
Date: Tue, 23 Nov 2021 13:08:27 GMT
Connection: keep-alive

# Not Frontend
curl: (28) Connection timed out after 5001 milliseconds
command terminated with exit code 28
```

Note that this is working despite the fact we did not delete the previous `backend-ingress-deny` policy:

- acend gmbh

```
kubectl get netpol
```

NAME	POD-SELECTOR	AGE
backend-allow-ingress-frontend	app=backend	2m7s
backend-ingress-deny	app=backend	12m

Network policies are additive. Just like with firewalls, it is thus a good idea to have default DENY policies and then add more specific ALLOW policies as needed.

10.2. Security contexts

In the concept of security context for a pod or container, there are several things to consider:

- Access control
- SELinux
- Running privileged or unprivileged workload
- Linux capabilities
- AppArmor
- Seccomp

In this lab you will learn where to configure and how to use some of these types.

Task 10.2.1: Access Control

Create a new pod by using this example:

```
apiVersion: v1
kind: Pod
metadata:
  name: security-context-demo
spec:
  securityContext:
    runAsUser: 1000
    runAsGroup: 3000
    fsGroup: 2000
  volumes:
  - name: sec-ctx-vol
    emptyDir: {}
  containers:
  - name: sec-ctx-demo
    image: busybox:1.28
    command: [ "sh", "-c", "sleep 1h" ]
    volumeMounts:
    - name: sec-ctx-vol
      mountPath: /data/demo
    securityContext:
      allowPrivilegeEscalation: false
```

You can see the different value entries in the 'securityContext' section, let's figure how what do they do. So create the pod and connect into the shell:

```
kubectl exec -it security-context-demo --namespace <namespace> -- sh
```

In the container run 'ps' to get a list of all running processes. The output shows, that the processes are running with the user 1000, which is the value from 'runAsUser':

PID	USER	TIME	COMMAND
1	1000	0:00	sleep 1h
6	1000	0:00	sh

Now navigate to the directory '/data' and list the content. As you can see the 'emptyDir' has been mounted with the group ID of 2000, which is the value of the 'fsGroup' field.

- acend gmbh

```
drwxrwsrwx 2 root 2000 4096 Oct 20 20:10 demo
```

Go into the dir 'demo' and create a file:

```
cd demo
echo hello > demofile
```

List the content with 'ls' again and see, that 'demofile' has the group ID 2000, which is the value 'fsGroup' as well.

Run the last command 'id' here and check the output:

```
uid=1000 gid=3000 groups=2000
```

The shown group ID of the user is 3000, from the field 'runAsGroup'. If the field would be empty the user would have 0 (root) and every process would be able to go with files which are owned by the root (0) group.

```
exit
```

Task 10.2.2: Advanced

As we are limited, in terms of permission, on the lab cluster we can't show all the other security contexts in a lab.

Check the documentation at kubernetes.io to view all the examples for [Security Contexts](#) .

11. Deployment strategies

In this lab, we are going to have a look at the different Deployment strategies.

This [document](#) should give you a good start. For more details, have a look at the [examples](#) or use [this demo](#) in which the different strategies are implemented as Helm charts.

Task 11.1: Apply deployment strategies

Apply some deployment strategies to the example from the [Scaling](#) .

12. Helm

[Helm](#) is a [Cloud Native Foundation](#) project to define, install and manage applications in Kubernetes.

tl;dr

Helm is a Package Manager for Kubernetes

- package multiple K8s resources into a single logical deployment unit
- ... but it's not just a Package Manager

Helm is a Deployment Management for Kubernetes

- do a repeatable deployment
- manage dependencies: reuse and share
- manage multiple configurations
- update, rollback and test application deployments

12.1. Helm overview

Ok, let's start with Helm. First, you have to understand the following 3 Helm concepts: **Chart**, **Repository** and **Release**.

A **Chart** is a Helm package. It contains all of the resource definitions necessary to run an application, tool, or service inside of a Kubernetes cluster. Think of it like the Kubernetes equivalent of a Homebrew formula, an Apt dpkg, or a Yum RPM file.

A **Repository** is the place where charts can be collected and shared. It's like Perl's CPAN archive or the Fedora Package Database, but for Kubernetes packages.

A **Release** is an instance of a chart running in a Kubernetes cluster. One chart can often be installed many times in the same cluster. Each time it is installed, a new release is created. Consider a MySQL chart. If you want two databases running in your cluster, you can install that chart twice. Each one will have its own release, which will in turn have its own release name.

With these concepts in mind, we can now explain Helm like this:

Helm installs charts into Kubernetes, creating a new release for each installation. To find new charts, you can search Helm chart repositories.

12.2. CLI installation

This guide shows you how to install the `helm` CLI tool. `helm` can be installed either from source or from pre-built binary releases. We are going to use the pre-built releases. `helm` binaries can be found on [Helm's release page](#) for the usual variety of operating systems.

Warning

If you do this training in our acend web based environment, no installation is required.

Task 12.2.1: Install CLI

Install the CLI for your **Operating System**

1. [Download the latest release](#)
2. Unpack it (e.g. `tar -zxvf <filename>`)
3. Copy to the correct location
 - Linux: Find the `helm` binary in the unpacked directory and move it to its desired destination (e.g. `mv linux-amd64/helm ~/.local/bin/`)
 - The desired destination should be listed in your `$PATH` environment variable (`echo $PATH`)
 - macOS: Find the `helm` binary in the unpacked directory and move it to its desired destination (e.g. `mv darwin-amd64/helm ~/bin/`)
 - The desired destination should be listed in your `$PATH` environment variable (`echo $PATH`)
 - Windows: Find the `helm` binary in the unpacked directory and move it to its desired destination
 - The desired destination should be listed in your `$PATH` environment variable (`echo $PATH`)

Task 12.2.2: Verify

To verify, run the following command and check if `Version` is what you expected:

```
helm version
```

The output is similar to this:

```
version.BuildInfo{Version:"v3.10.1", GitCommit:"9f88ccb6aee40b9a0535fcc7efea6055e1ef72c9", GitTreeState:"clean", GoVersion:"go1.18.7"}
```

From here on you should be able to run the client.

12.3. Create a chart

In this lab we are going to create our very first Helm chart and deploy it.

Task 12.3.1: Create Chart

First, let's create our chart. Open your favorite terminal and make sure you're in the workspace for this lab, e.g. `cd ~/<workspace-kubernetes-training>` :

```
helm create mychart
```

You will now find a `mychart` directory with the newly created chart. It already is a valid and fully functional chart which deploys an nginx instance. Have a look at the generated files and their content. For an explanation of the files, visit the [Helm Developer Documentation](#) . In a later section you'll find all the information about Helm templates.

Task 12.3.2: Install Release

Before actually deploying our generated chart, we can check the (to be) generated Kubernetes resources with the following command:

```
helm install --dry-run --debug --namespace <namespace> myfirstrelease ./mychart
```

Finally, the following command creates a new release and deploys the application:

```
helm install --namespace <namespace> myfirstrelease ./mychart
```

With `kubectl get pods --namespace <namespace>` you should see a new Pod:

NAME	READY	STATUS	RESTARTS	AGE
myfirstrelease-mychart-6d4956b75-ng8x4	1/1	Running	0	2m21s

You can list the newly created Helm release with the following command:

```
helm ls --namespace <namespace>
```

Task 12.3.3: Expose Application

Our freshly deployed nginx is not yet accessible from outside the Kubernetes cluster. To expose it, we have to make sure a so called ingress resource will be deployed as well.

- acend gmbh

Also make sure the application is accessible via TLS.

A look into the file `templates/ingress.yaml` reveals that the rendering of the ingress and its values is configurable through `values( values.yaml )`:

```
{{- if .Values.ingress.enabled -}}
{{- $fullName := include "mychart.fullname" . -}}
{{- $svcPort := .Values.service.port -}}
{{- if and .Values.ingress.className (not (semverCompare ">=1.18-0" .Capabilities.KubeVersion.GitVersion)) }}
  {{- if not (hasKey .Values.ingress.annotations "kubernetes.io/ingress.class") }}
    {{- $_ := set .Values.ingress.annotations "kubernetes.io/ingress.class" .Values.ingress.className }}
  {{- end }}
{{- end }}
{{- if semverCompare ">=1.19-0" .Capabilities.KubeVersion.GitVersion -}}
apiVersion: networking.k8s.io/v1
{{- else if semverCompare ">=1.14-0" .Capabilities.KubeVersion.GitVersion -}}
apiVersion: networking.k8s.io/v1beta1
{{- else -}}
apiVersion: extensions/v1beta1
{{- end }}
kind: Ingress
metadata:
  name: {{ $fullName }}
  labels:
    {{- include "mychart.labels" . | nindent 4 }}
    {{- with .Values.ingress.annotations }}
  annotations:
    {{- toYaml . | nindent 4 }}
    {{- end }}
spec:
  {{- if and .Values.ingress.className (semverCompare ">=1.18-0" .Capabilities.KubeVersion.GitVersion) }}
  ingressClassName: {{ .Values.ingress.className }}
  {{- end }}
  {{- if .Values.ingress.tls }}
  tls:
    {{- range .Values.ingress.tls }}
    - hosts:
        {{- range .hosts }}
        - {{ . | quote }}
        {{- end }}
      secretName: {{ .secretName }}
    {{- end }}
  {{- end }}
  rules:
    {{- range .Values.ingress.hosts }}
    - host: {{ .host | quote }}
      http:
        paths:
          {{- range .paths }}
          - path: {{ .path }}
            {{- if and .pathType (semverCompare ">=1.18-0" $.Capabilities.KubeVersion.GitVersion) }}
            pathType: {{ .pathType }}
            {{- end }}
          backend:
            {{- if semverCompare ">=1.19-0" $.Capabilities.KubeVersion.GitVersion }}
            service:
              name: {{ $fullName }}
              port:
                number: {{ $svcPort }}
            {{- else }}
            serviceName: {{ $fullName }}
            servicePort: {{ $svcPort }}
            {{- end }}
          {{- end }}
    {{- end }}
  {{- end }}
```

Thus, we need to change this value inside our `mychart/values.yaml` file. This is also where we enable the TLS part:

Note

Make sure to replace the `<namespace>` and `<appdomain>` accordingly.

```
[...]
ingress:
  enabled: true
  hosts:
    - host: mychart-<namespace>.<appdomain>
      paths:
        - path: /
          pathType: ImplementationSpecific
[...]
```

Note

Make sure to set the proper value as hostname. `<appdomain>` will be provided by the trainer.

Apply the change by upgrading our release:

```
helm upgrade --namespace <namespace> myfirstrelease ./mychart
```

This will result in something similar to:

```
Release "myfirstrelease" has been upgraded. Happy Helming!
NAME: myfirstrelease
LAST DEPLOYED: Wed Dec  2 14:44:42 2020
NAMESPACE: <namespace>
STATUS: deployed
REVISION: 2
NOTES:
1. Get the application URL by running these commands:
  https://<namespace>.<appdomain>/
```

Check whether the ingress was successfully deployed by accessing the URL `https://mychart-<namespace>.<appdomain>/`

Task 12.3.4: Overwrite value using commandline param

An alternative way to set or overwrite values for charts we want to deploy is the `--set name=value` parameter. This parameter can be used when installing a chart as well as upgrading.

Update the replica count of your nginx Deployment to 2 using `--set name=value`

Solution

```
helm upgrade --namespace <namespace> --set replicaCount=2 myfirstrelease ./mychart
```

Values that have been set using `--set` can be reset by helm upgrade with `--reset-values`.

Task 12.3.5: Values

Have a look at the `values.yaml` file in your chart and study all the possible configuration params introduced in a freshly created chart.

Task 12.3.6: Remove release

To remove an application, simply remove the Helm release with the following command:

```
helm uninstall myfirstrelease --namespace <namespace>
```

Do this with our deployed release. With `kubect1 get pods --namespace <namespace>` you should no longer see your application Pod.

12.4. Complex example

In this extended lab, we are going to deploy an existing, more complex application with a Helm chart from the Artifact Hub.

Artifact Hub

Check out [Artifact Hub](#) where you'll find a huge number of different Helm charts. For this lab, we'll use the [WordPress chart by Bitnami](#), a publishing platform for building blogs and websites.

WordPress

As this WordPress Helm chart is published in Bitnami's Helm repository, we're first going to add it to our local repo list:

```
helm repo add bitnami https://charts.bitnami.com/bitnami
```

Let's check if that worked:

```
helm repo list
```

```
NAME    URL
bitnami https://charts.bitnami.com/bitnami
```

Now look at the available configuration for this Helm chart. Usually you can find it in the [values.yaml](#) or in the chart's readme file. You can also check it on its [Artifact Hub page](#).

We are going to override some of the values. For that purpose, create a new `values.yaml` file locally on your workstation (e.g. `~/<workspace>/values.yaml`) with the following content:

```
---
persistence:
  size: 1Gi
service:
  type: ClusterIP
updateStrategy:
  type: Recreate

ingress:
  enabled: true
  hostname: wordpress-<namespace>.<appdomain>
  extraTls:
  - hosts:
    - wordpress-<namespace>.<appdomain>

mariadb:
  primary:
    persistence:
      size: 1Gi
```

Note

Make sure to set the proper value as hostname. `<appdomain>` will be provided by the trainer.

If you look inside the [Chart.yaml](#) file of the WordPress chart, you'll see a dependency to the [MariaDB Helm chart](#). All the MariaDB values are used by this dependent Helm chart and the chart is automatically deployed when installing WordPress.

The `Chart.yaml` file allows us to define dependencies on other charts. In our Wordpress chart we use the `Chart.yaml` to add a `mariadb` to store the WordPress data in.

```
dependencies:
- condition: mariadb.enabled
  name: mariadb
  repository: https://charts.bitnami.com/bitnami
  version: 9.x.x
```

[Helm's best practices](#) suggest to use version ranges instead of a fixed version whenever possible. The suggested default therefore is patch-level version match:

```
version: ~3.5.7
```

This is e.g. equivalent to `>= 3.5.7, < 3.6.0`. Check [this SemVer readme chapter](#) for more information on version ranges.

Note

For more details on how to manage **dependencies**, check out the [Helm Dependencies Documentation](#).

Subcharts are an alternative way to define dependencies within a chart: A chart may contain another chart (inside of its `charts/` directory) upon which it depends. As a result, when installing the chart, it will install all of its dependencies from the `charts/` directory.

We are now going to deploy the application in a specific version (which is not the latest release on purpose). Also note that we define our custom `values.yaml` file with the `-f` parameter:

```
helm install wordpress bitnami/wordpress -f values.yaml --namespace <namespace>
```

Look for the newly created resources with `helm ls` and `kubectl get deploy,pod,ingress,pvc`:

```
helm ls --namespace <namespace>
```

which gives you:

NAME	NAMESPACE	REVISION	UPDATED	STATUS	CHART	APP VERSION
wordpress	<namespace>	1	2021-03-25 14:27:38.231722961 +0100 CET	deployed	wordpress-10.7.1	5.7.0

- acend gmbh

and

```
kubectl get deploy,pod,ingress,pvc --namespace <namespace>
```

which gives you:

NAME	READY	UP-TO-DATE	AVAILABLE	AGE
deployment.apps/wordpress	1/1	1	1	2m6s

NAME	READY	STATUS	RESTARTS	AGE
pod/wordpress-6bf6df9c5d-w4fpx	1/1	Running	0	2m6s
pod/wordpress-mariadb-0	1/1	Running	0	2m6s

NAME	HOSTS	ADDRESS	PORTS	AGE
ingress.extensions/wordpress	wordpress-<namespace>.<appdomain>	10.100.1.10	80	2m6s

NAME	MODES	STORAGECLASS	AGE	STATUS	VOLUME	CAPACITY	ACCESS
persistentvolumeclaim/data-wordpress-mariadb-0		cloudscale-volume-ssd	2m6s	Bound	pvc-859fe3b4-b598-4f86-b7ed-a3a183f700fd	1Gi	RWO
persistentvolumeclaim/wordpress		cloudscale-volume-ssd	2m7s	Bound	pvc-83ebf739-0b0e-45a2-936e-e925141a0d35	1Gi	RWO

In order to check the values used in a given release, execute:

```
helm get values wordpress --namespace <namespace>
```

which gives you:

```
USER-SUPPLIED VALUES:
ingress:
  enabled: true
  hostname: wordpress-<namespace>.<appdomain>
mariadb:
  primary:
    persistence:
      size: 1Gi
persistence:
  size: 1Gi
service:
  type: ClusterIP
updateStrategy:
  type: Recreate
```

As soon as all deployments are ready (meaning pods `wordpress` and `mariadb` are running) you can open the application with the URL from your Ingress resource defined in `values.yaml`.

Upgrade

We are now going to upgrade the application to a newer Helm chart version. When we installed the Chart, a couple of secrets were needed during this process. In order to do the upgrade of the Chart now, we need to provide those secrets to the upgrade command, to be sure no sensitive data will be overwritten:

- acend gmbh

- wordpressPassword
- mariadb.auth.rootPassword
- mariadb.auth.password

Note

This is specific to the wordpress Bitami Chart, and might be different when installing other Charts.

Use the following commands to gather the secrets and store them in environment variables. Make sure to replace `<namespace>` with your current value.

```
export WORDPRESS_PASSWORD=$(kubectl get secret wordpress -o jsonpath="{.data.wordpress-password}" --namespace <namespace> | base64 --decode)
```

```
export MARIADB_ROOT_PASSWORD=$(kubectl get secret wordpress-mariadb -o jsonpath="{.data.mariadb-root-password}" --namespace <namespace> | base64 --decode)
```

```
export MARIADB_PASSWORD=$(kubectl get secret wordpress-mariadb -o jsonpath="{.data.mariadb-password}" --namespace <namespace> | base64 --decode)
```

Then do the upgrade with the following command:

```
helm upgrade -f values.yaml --set wordpressPassword=$WORDPRESS_PASSWORD --set mariadb.auth.rootPassword=$MARIADB_ROOT_PASSWORD --set mariadb.auth.password=$MARIADB_PASSWORD wordpress bitnami/wordpress --namespace <namespace>
```

And then observe the changes in your WordPress and MariaDB Apps

Cleanup

```
helm uninstall wordpress --namespace <namespace>
```

Additional Task

Study the Helm [best practices](#) as an optional and additional task.

13. Kustomize

[Kustomize](#) is a tool to manage YAML configurations for Kubernetes objects in a declarative and reusable manner. In this lab, we will use Kustomize to deploy the same app for two different environments.

Installation

Kustomize can be used in two different ways:

- As a standalone `kustomize` binary, downloadable from [here](#)
- With the parameter `--kustomize` or `-k` in certain `kubectl` subcommands such as `apply` or `create`

Note

You might get a different behaviour depending on which variant you use. The reason for this is that the version built into `kubectl` is usually older than the standalone binary.

Usage

The main purpose of Kustomize is to build configurations from a predefined file structure (which will be introduced in the next section):

```
kustomize build <dir>
```

The same can be achieved with `kubectl` :

```
kubectl kustomize <dir>
```

The next step is to apply this configuration to the Kubernetes cluster:

```
kustomize build <dir> | kubectl apply -f -
```

Or in one `kubectl` command with the parameter `-k` instead of `-f` :

```
kubectl apply -k <dir>
```

Task 13.1: Prepare a Kustomize config

We are going to deploy a simple application:

- The Deployment starts an application based on nginx
- A Service exposes the Deployment

- acend gmbh

- The application will be deployed for two different example environments, integration and production

Kustomize allows inheriting Kubernetes configurations. We are going to use this to create a base configuration and then override it for the different environments. Note that Kustomize does not use templating. Instead, smart patch and extension mechanisms are used on plain YAML manifests to keep things as simple as possible.

Get the example config

Find the needed resource files inside the folder `content/en/docs/kustomize/kustomize` of the techlab github repository. Clone the [repository](#) or get the content as [zip](#)

Change to the folder `content/en/docs/kustomize/kustomize` to execute the kustomize commands.

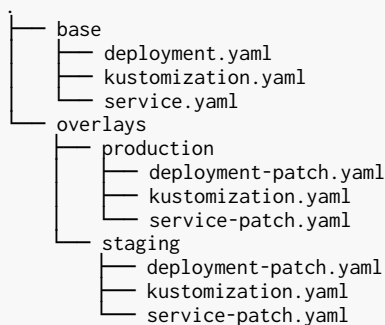
Note

Commands for git checkout and folder switch:

```
git clone https://github.com/acend/kubernetes-basics-training.git
cd kubernetes-basics-training/content/en/docs/kustomize/kustomize/
```

File structure

The structure of a Kustomize configuration typically looks like this:



Base

Let's have a look at the `base` directory first which contains the base configuration. There's a `deployment.yaml` with the following content:

- acend gmbh

```
apiVersion: apps/v1
kind: Deployment
metadata:
  name: kustomize-app
spec:
  selector:
    matchLabels:
      app: kustomize-app
  template:
    metadata:
      labels:
        app: kustomize-app
    spec:
      containers:
        - name: kustomize-app
          image: quay.io/acend/example-web-go
          env:
            - name: APPLICATION_NAME
              value: app-base
          command:
            - sh
            - -c
            - |-
              set -e
              /bin/echo "My name is $APPLICATION_NAME"
              /usr/local/bin/go
          ports:
            - name: http
              containerPort: 80
              protocol: TCP
```

There's also a Service for our Deployment in the corresponding `base/service.yaml` :

```
apiVersion: v1
kind: Service
metadata:
  name: kustomize-app
spec:
  ports:
    - port: 80
      targetPort: 80
  selector:
    app: kustomize-app
```

And there's an additional `base/kustomization.yaml` which is used to configure Kustomize:

```
resources:
  - service.yaml
  - deployment.yaml
```

It references the previous manifests `service.yaml` and `deployment.yaml` and makes them part of our base configuration.

Overlays

Now let's have a look at the other directory which is called `overlays` . It contains two subdirectories `staging` and `production` which both contain a `kustomization.yaml` with almost the same content.

`overlays/staging/kustomization.yaml` :

- acend gmbh

```
nameSuffix: -staging
bases:
- ../../base
patchesStrategicMerge:
- deployment-patch.yaml
- service-patch.yaml
```

overlays/production/kustomization.yaml :

```
nameSuffix: -production
bases:
- ../../base
patchesStrategicMerge:
- deployment-patch.yaml
- service-patch.yaml
```

Only the first key `nameSuffix` differs.

In both cases, the `kustomization.yaml` references our base configuration. However, the two directories contain two different `deployment-patch.yaml` files which patch the `deployment.yaml` from our base configuration.

overlays/staging/deployment-patch.yaml :

```
apiVersion: apps/v1
kind: Deployment
metadata:
  name: kustomize-app
spec:
  selector:
    matchLabels:
      app: kustomize-app-staging
  template:
    metadata:
      labels:
        app: kustomize-app-staging
    spec:
      containers:
        - name: kustomize-app
          env:
            - name: APPLICATION_NAME
              value: kustomize-app-staging
```

overlays/production/deployment-patch.yaml :

- acend gmbh

```
apiVersion: apps/v1
kind: Deployment
metadata:
  name: kustomize-app
spec:
  selector:
    matchLabels:
      app: kustomize-app-production
  template:
    metadata:
      labels:
        app: kustomize-app-production
    spec:
      containers:
        - name: kustomize-app
          env:
            - name: APPLICATION_NAME
              value: kustomize-app-production
```

The main difference here is that the environment variable `APPLICATION_NAME` is set differently. The `app` label also differs because we are going to deploy both Deployments into the same Namespace.

The same applies to our Service. It also comes in two customizations so that it matches the corresponding Deployment in the same Namespace.

overlays/staging/service-patch.yaml :

```
apiVersion: v1
kind: Service
metadata:
  name: kustomize-app
spec:
  selector:
    app: kustomize-app-staging
```

overlays/production/service-patch.yaml :

```
apiVersion: v1
kind: Service
metadata:
  name: kustomize-app
spec:
  selector:
    app: kustomize-app-production
```

Note

All files mentioned above are also directly accessible from [GitHub](#) .

Prepare the files as described above in a local directory of your choice.

Task 13.2: Deploy with Kustomize

We are now ready to deploy both apps for the two different environments. For simplicity, we will use the same Namespace.

- acend gmbh

```
kubectl apply -k overlays/staging --namespace <namespace>
```

```
service/kustomize-app-staging created  
deployment.apps/kustomize-app-staging created
```

```
kubectl apply -k overlays/production --namespace <namespace>
```

```
service/kustomize-app-production created  
deployment.apps/kustomize-app-production created
```

As you can see, we now have two deployments and services deployed. Both of them use the same base configuration. However, they have a specific configuration on their own as well.

Let's verify this. Our app writes a corresponding log entry that we can use for analysis:

```
kubectl get pods --namespace <namespace>
```

NAME	READY	STATUS	RESTARTS	AGE
kustomize-app-production-74c7bdb7d-8cccd	1/1	Running	0	2m1s
kustomize-app-staging-7967885d5b-qp6l8	1/1	Running	0	5m33s

```
kubectl logs kustomize-app-staging-7967885d5b-qp6l8
```

```
My name is kustomize-app-staging
```

```
kubectl logs kustomize-app-production-74c7bdb7d-8cccd
```

```
My name is kustomize-app-production
```

Further information

Kustomize has more features of which we just covered a couple. Please refer to the docs for more information.

- Kustomize documentation: <https://kubernetes-sigs.github.io/kustomize/>
- API reference: <https://kubernetes-sigs.github.io/kustomize/api-reference/>
- Another `kustomization.yaml` reference: <https://kubectl.docs.kubernetes.io/pages/reference/kustomize.html>

- acend gmbh

- Examples: <https://github.com/kubernetes-sigs/kustomize/tree/master/examples>