

Autoscaling in Kubernetes

From Zero to Hero



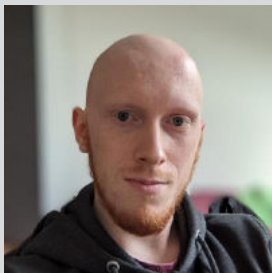
Lukas Paluch Tobias Manske

12. März 2023

Hello



Wer sind wir



mail manske@atix.de
in [linkedin.com/in/tmanske](https://www.linkedin.com/in/tmanske)
gh github.com/rad4day

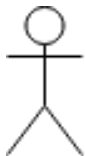


mail paluch@atix.de
in [linkedin.com/in/lukasp](https://www.linkedin.com/in/lukasp)
gh github.com/fluktuid

ATIX AG

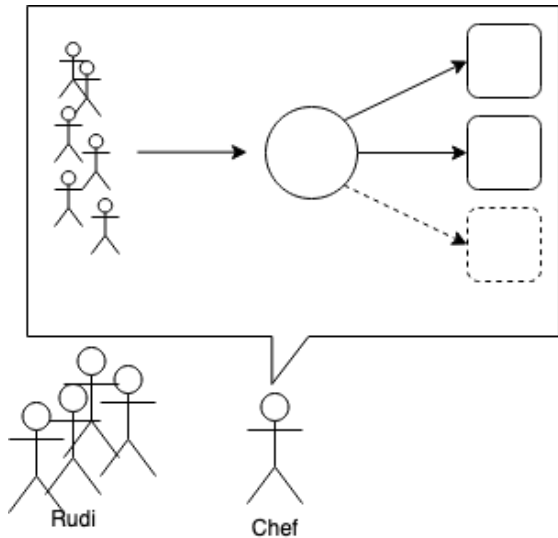
- ▶ Transformation und Automatisierung von Rechenzentren
- ▶ Fokus auf
 - ▶ Lifecycle System Management
 - ▶ Container-Technologie

Whats the Story?



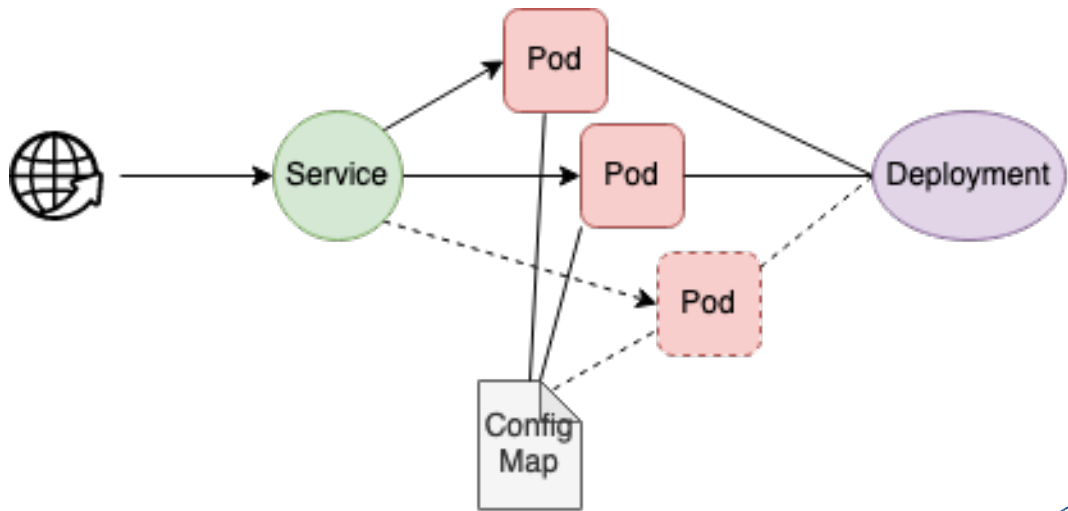
Rudi

Whats the Story?



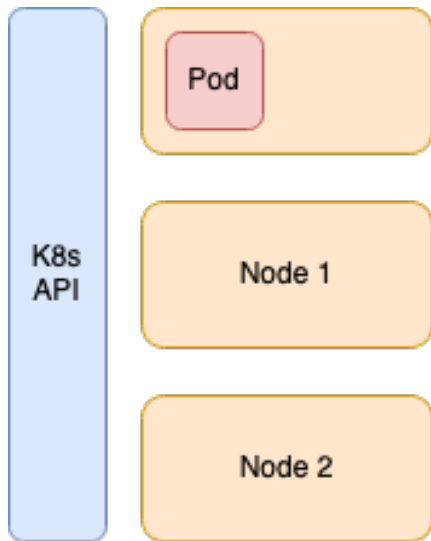
- 1 Whats Kubernetes
- 2 Setup
- 3 Skalierung
 - Skalierung des Clusters
 - Scale to Zero

- ▶ Verteiltes System
- ▶ Betrieb (Orchestrierung) von Containern
- ▶ Gut skalierbar
- ▶ Möglichkeit zu HA-Setup

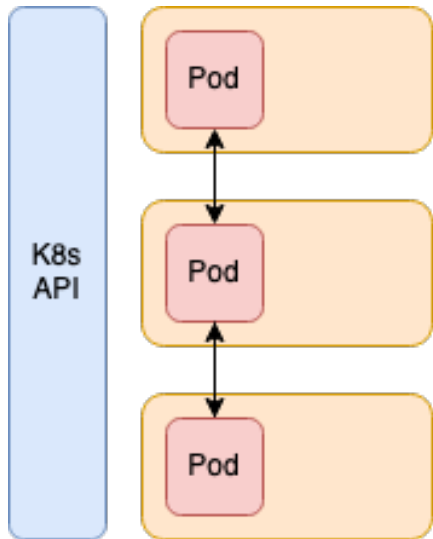


- 1 Whats Kubernetes
- 2 Setup
- 3 Skalierung
 - Skalierung des Clusters
 - Scale to Zero

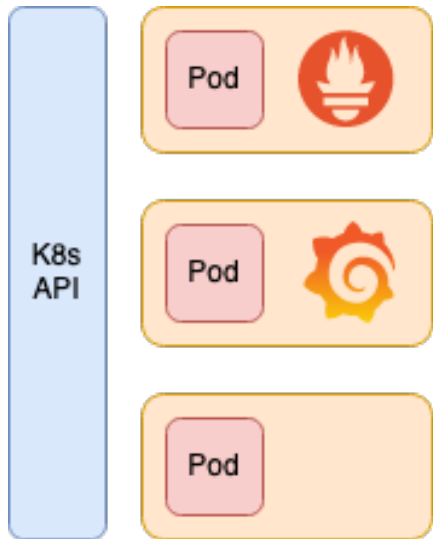
Setup



Setup Replizierung + Affinity



Setup Monitoring

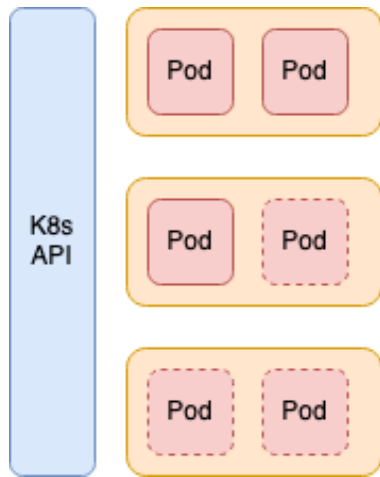
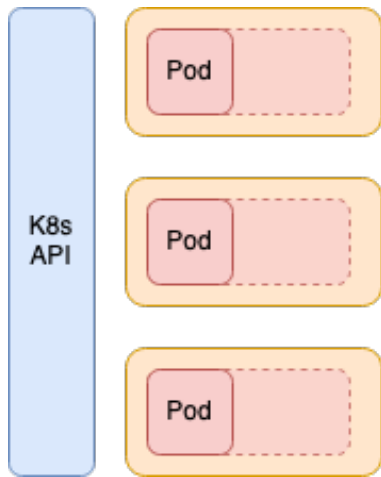


- 1 Whats Kubernetes
- 2 Setup
- 3 Skalierung**
 - Skalierung des Clusters
 - Scale to Zero

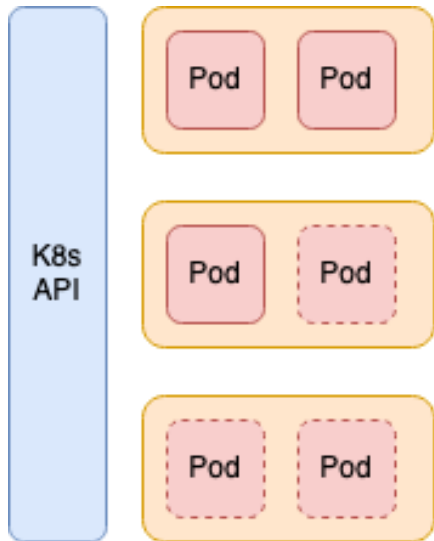
Skalierung der Anwendung



PodAutoscaler

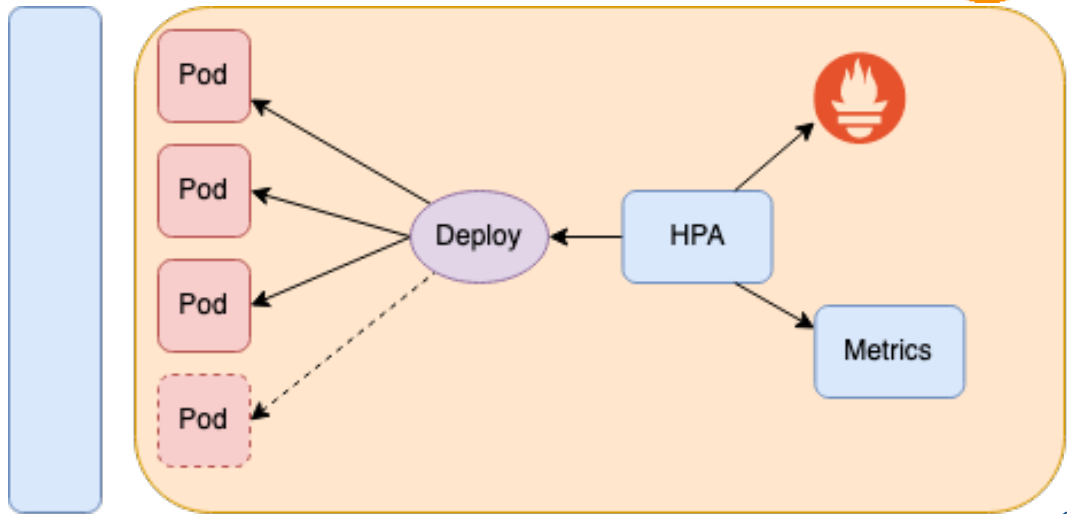


HorizontalPodAutoscaler



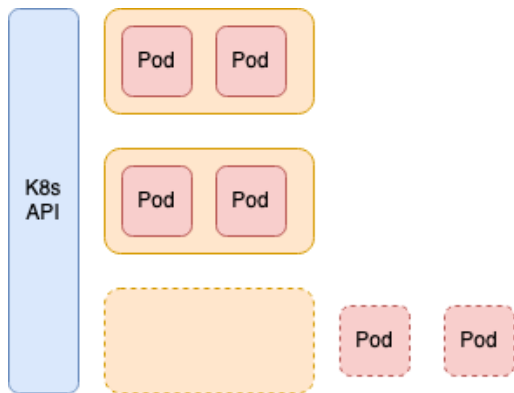
```
1 kind: HorizontalPodAutoscaler
2 ...
3 spec:
4   scaleTargetRef:
5     apiVersion: apps/v1
6     kind: Deployment
7     name: server
8   minReplicas: 1
9   maxReplicas: 100
10  metrics:
11    - type: Pods
12      pods:
13        metric:
14          name:
15                                server_requests_per_second
16
17    target:
18      type: AverageValue
19      averageValue: 16
20  ...
```

So HPA

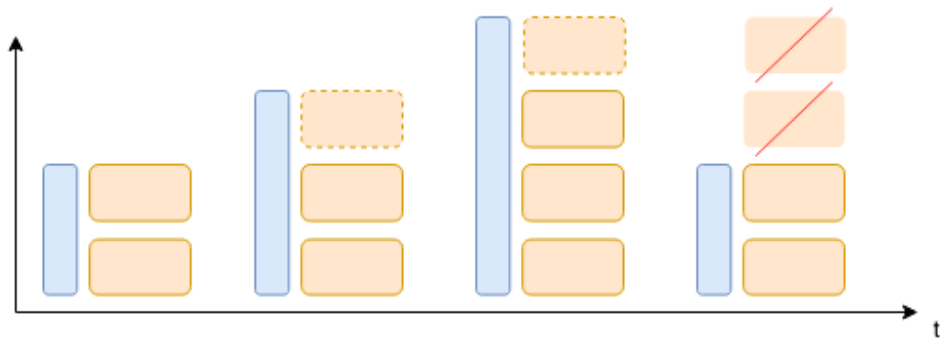


- 1 Whats Kubernetes
- 2 Setup
- 3 Skalierung
 - Skalierung des Clusters
 - Scale to Zero

Skalierung des Clusters



```
1 # cluster-autoscaler.yaml
2 autoscalingGroups:
3   - name: "<eks-worker-group>"
4     minSize: 1
5     maxSize: 20
```



- 1 Whats Kubernetes
- 2 Setup
- 3 Skalierung
 - Skalierung des Clusters
 - Scale to Zero

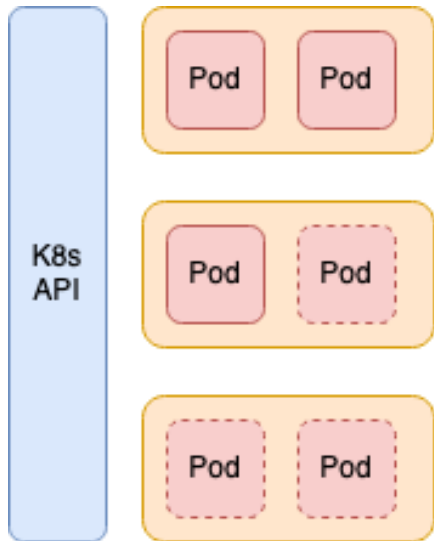
X auf 0

- ▶ Einfach umzusetzen.
- ▶ Anwendung »verschwindet« wenn keine Last vorhanden.

0 auf X

- ▶ Benötigt podunabhängige Metrik.
- ▶ Was passiert mit Anfragen wenn kein Pod vorhanden?
- ▶ (Knative vs KEDA?)

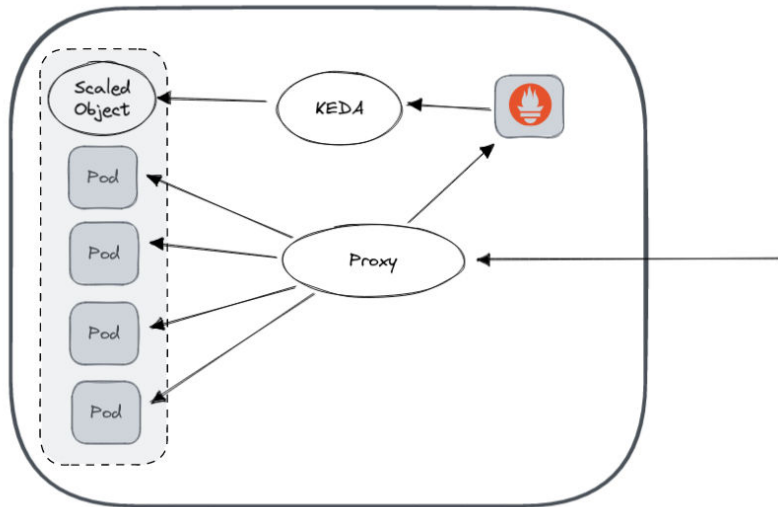
Scale to Zero

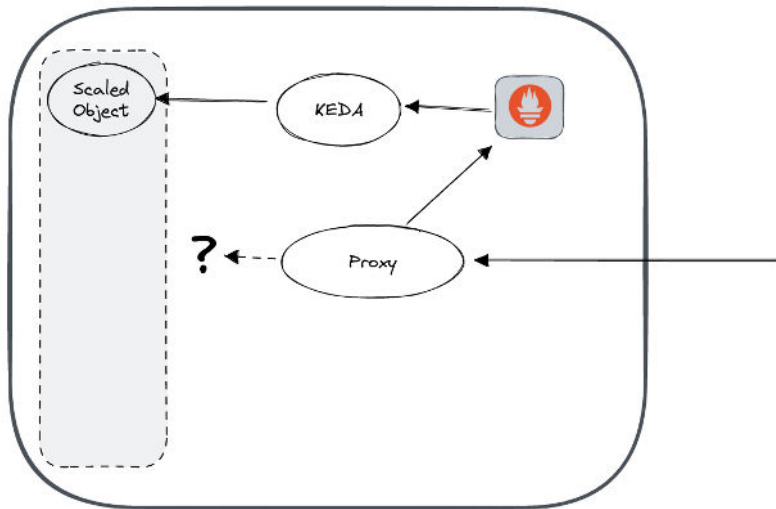


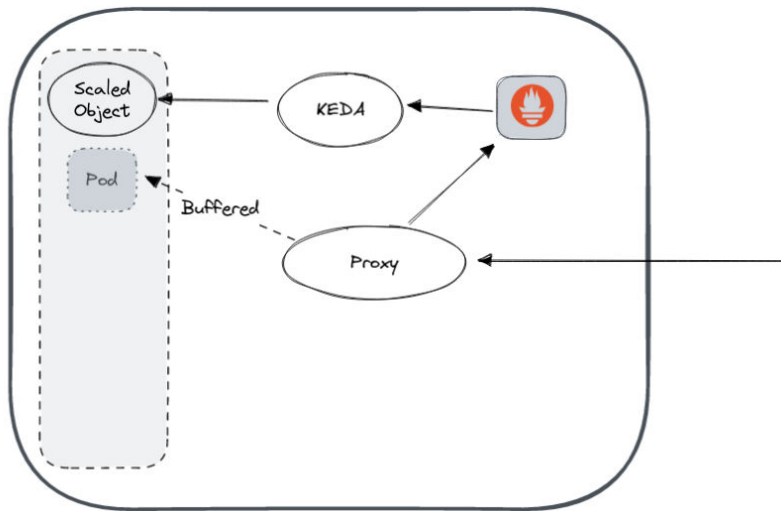
```
1 apiVersion: keda.sh/v1alpha1
2 kind: ScaledObject
3 metadata:
4   ...
5 spec:
6   scaleTargetRef:
7     name: scaled-deployment
8   pollingInterval: 10
9   cooldownPeriod: 15
10  minReplicaCount: 0 # <--
11  maxReplicaCount: 10
12  triggers:
13    - type: prometheus
14      metadata:
15        serverAddress: http://...
16        metricName: access_frequency
17        threshold: '25'
18        query: sum(rate(
19          node_http_requests_total
20          [2m]))
```


Aus dem keda.sh Blog

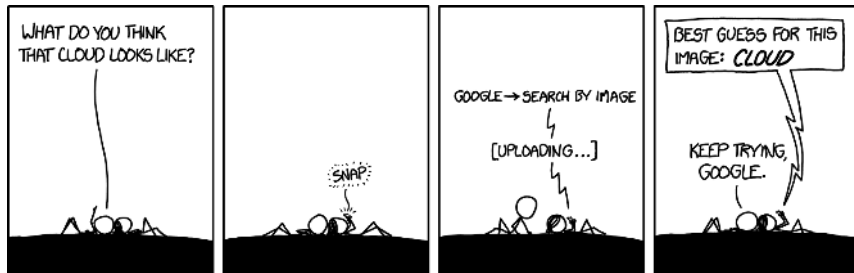
Autoscaling HTTP is often not as straightforward as other event sources. You don't know how much traffic will be coming and, given its synchronous nature, supporting scale-to-zero HTTP applications requires an intelligent intermediate routing layer to "hold" the incoming request(s) until new instances of the backend application are created and running.







And now there is Time for you



(xkcd.com/1444/)

presentation & code
github.com/fluktuid/talk-k8s-scaling