



CLOUD-NATIVE & KUBERNETES

Kubernetes en Production

La Checklist Complète

Faire tourner un cluster ne suffit pas. Ce guide détaille les contrôles de disponibilité, de sécurité et d'observabilité qui séparent un cluster de développement d'une plateforme réellement prête pour la production.

11 pages · Checklist d'auto-évaluation, comparatif RBAC/réseau/conteneurs et feuille de route incluses

82%

des utilisateurs de conteneurs font tourner Kubernetes en production (CNCF, 2026)

36%

citent la sécurité comme frein principal à l'adoption (CNCF)

25%

des incidents cloud impliquent une mauvaise configuration (IBM X-Force)

SOMMAIRE

Ce que vous allez trouver dans ce document

Un parcours en six temps, de la définition de la « vraie » mise en production à une feuille de route opérationnelle concrète.

01	Synthèse — l'essentiel en une page	p.3
02	Pourquoi la mise en production est le vrai test	p.4
03	Disponibilité : haute dispo, mises à jour, arrêts propres	p.5
04	Sécurité : RBAC, réseau, conteneurs	p.6
05	Observabilité et santé des applications	p.7
06	Ressources, stockage et autoscaling	p.8
07	Grille d'auto-évaluation « prêt pour la prod »	p.9
08	Feuille de route en 5 étapes	p.10
09	Conclusion & à propos de Move2Cloud	p.11

À qui s'adresse ce document ? Équipes plateforme, SRE, DevOps et RSSI qui préparent — ou consolident — un cluster Kubernetes destiné à accueillir des charges critiques.

SYNTHÈSE

L'essentiel en une page

Kubernetes est devenu le standard de fait pour faire tourner des charges en production — mais « ça tourne » et « c'est prêt pour la production » sont deux choses différentes.

Selon l'édition 2026 de l'enquête annuelle CNCF, **82 % des utilisateurs de conteneurs font désormais tourner Kubernetes en production**, contre 66 % en 2023. L'adoption technique n'est plus le principal obstacle : **47 % des organisations citent les changements culturels côté équipes de développement** comme premier frein, devant la formation (36 %) et la sécurité (36 %).

Cette bascule culturelle a un revers opérationnel : de nombreux clusters passent en production avec des lacunes qui restent invisibles tant que rien ne casse — absence de PodDisruptionBudget, RBAC trop permissif, sondes de santé mal calibrées, ou requêtes/limites de ressources non définies. Ce sont précisément ces lacunes que la mauvaise configuration — et non l'exploitation d'une vulnérabilité — transforme en incident : **25 % des incidents de sécurité cloud impliquent un service mal configuré**, selon l'IBM X-Force Threat Intelligence Index.

82%

des utilisateurs de conteneurs exploitent Kubernetes en production en 2026 (CNCF)

47%

citent la culture d'équipe — pas la technique — comme premier frein à l'adoption

25%

des incidents cloud proviennent d'une configuration mal maîtrisée

Pour une équipe qui prépare une mise en production — premier cluster critique ou consolidation d'un existant — la question n'est plus « Kubernetes est-il le bon choix ? » mais « quels contrôles de disponibilité, de sécurité et d'observabilité sont réellement en place avant d'y faire tourner du trafic réel ? ». Ce document propose une checklist structurée pour y répondre.

SECTION 02

Pourquoi la mise en production est le vrai test

Un cluster Kubernetes qui fonctionne en environnement de développement valide surtout une chose : que les manifestes s'appliquent sans erreur. Il ne dit rien de sa capacité à absorber une mise à jour de nœud, une panne de zone de disponibilité, un pic de trafic ou une tentative d'accès non autorisé. Trois écarts expliquent pourquoi tant d'incidents de production trouvent leur origine dans des choix faits — ou non faits — bien avant le premier déploiement réel.

1. Le développement masque les défaillances partielles

En développement, un pod qui redémarre ou une requête qui échoue une fois passe souvent inaperçu. En production, ces mêmes symptômes signalent une absence de réplication suffisante, de sondes de santé correctement calibrées, ou de gestion propre des arrêts — des lacunes qui ne se révèlent qu'sous charge réelle.

2. La sécurité par défaut de Kubernetes reste permissive

Les paramètres par défaut de Kubernetes privilégient la simplicité de démarrage, pas la sécurité en production : RBAC ouvert, absence de politiques réseau, conteneurs exécutés en root si rien ne l'interdit. Rien n'empêche un cluster de fonctionner parfaitement tout en étant largement exposé.

3. La complexité opérationnelle grandit avec l'échelle

Un cluster à trois services est simple à observer à l'œil nu. Un cluster à cent services ne l'est plus : sans observabilité structurée (métriques, logs, traces) et sans automatisation du dimensionnement, les équipes perdent la visibilité nécessaire pour détecter une dérive avant qu'elle ne devienne un incident.

« Un cluster Kubernetes n'est pas prêt pour la production parce qu'il tourne — il l'est parce qu'il a été conçu pour survivre à ce qui va mal tourner. »

Les sections suivantes détaillent, domaine par domaine, les contrôles concrets qui comblent ces trois écarts.

SECTION 03

Disponibilité : haute dispo, mises à jour, arrêts propres

La disponibilité ne se décrète pas — elle se construit à travers un ensemble de réglages qui déterminent comment le cluster réagit face à une panne ou une mise à jour.

CONTRÔLE	CE QU'IL FAUT METTRE EN PLACE
Réplication multi-zone	Déployer au moins 3 réplicas par service critique, répartis sur plusieurs zones de disponibilité, avec des règles d'anti-affinité pour éviter que tous les pods d'un même service finissent sur le même nœud.
PodDisruptionBudget	Définir un budget de disruption pour chaque charge critique, afin d'empêcher qu'une mise à jour de nœud ou une opération de maintenance ne fasse tomber trop de réplicas simultanément.
Mises à jour progressives	Utiliser des stratégies de déploiement progressif (rolling update, blue/green ou canary) plutôt que des remplacements brutaux, avec des seuils maxUnavailable/maxSurge explicitement définis.
Arrêt propre (graceful shutdown)	Gérer correctement le signal SIGTERM applicatif et définir une terminationGracePeriodSeconds cohérente avec le temps réel de fermeture des connexions en cours.

À retenir : une haute disponibilité bien conçue se vérifie en la testant — en coupant volontairement un nœud ou une zone en environnement de pré-production — plutôt qu'en la supposant acquise parce que le manifeste déclare plusieurs réplicas.

Le piège du réplica unique

Un service déployé avec un seul réplica n'a, par définition, aucune tolérance aux pannes : toute défaillance de pod, toute mise à jour de nœud, tout redémarrage planifié entraîne une interruption de service. C'est l'écart de configuration le plus simple à corriger, et l'un des plus fréquents dans les clusters qui n'ont pas encore subi d'audit de mise en production.

SECTION 04

Sécurité : RBAC, réseau, conteneurs

La sécurité Kubernetes se joue à trois niveaux distincts — qui peut agir sur le cluster, qui peut communiquer avec qui, et ce qu'un conteneur compromis peut atteindre.

NIVEAU	CONTRÔLE ATTENDU
RBAC	Appliquer le principe du moindre privilège : rôles scopés par espace de noms plutôt que des ClusterRoleBinding larges, comptes de service dédiés par application, revue régulière des permissions accordées.
Politiques réseau	Définir des NetworkPolicy explicites plutôt que de laisser le trafic est-ouest totalement ouvert par défaut — un pod compromis ne doit pas pouvoir atteindre librement le reste du cluster.
Durcissement des conteneurs	Exécuter les conteneurs en utilisateur non-root, système de fichiers racine en lecture seule quand c'est possible, capacités Linux réduites au strict nécessaire, images scannées avant déploiement.
Secrets	Ne jamais stocker de secrets en clair dans les manifestes ; utiliser le chiffrement au repos d'etcd et, idéalement, un gestionnaire de secrets externe intégré au cluster.

Point de vigilance : selon l'IBM X-Force Threat Intelligence Index, 25 % des incidents de sécurité cloud proviennent d'une configuration mal maîtrisée — pas d'une vulnérabilité inédite. Les quatre contrôles ci-dessus couvrent la grande majorité des angles d'attaque réellement exploités contre des clusters Kubernetes.

Le tableau de bord Kubernetes, une porte trop souvent ouverte

Un tableau de bord ou une API server exposé sans authentification forte ni restriction réseau reste l'une des erreurs de configuration les plus documentées sur Kubernetes. La règle simple : aucune interface d'administration ne devrait être accessible sans authentification, quelle que soit la sensibilité perçue du cluster.

SECTION 05

Observabilité et santé des applications

Sans observabilité structurée, un cluster à l'échelle devient impossible à diagnostiquer autrement qu'en réaction à un incident déjà en cours.

Des sondes de santé correctement calibrées

Les sondes *liveness*, *readiness* et *startup* ne sont pas interchangeables : la *liveness* redémarre un pod bloqué, la *readiness* le retire du service tant qu'il n'est pas prêt à recevoir du trafic, la *startup* protège les applications à démarrage lent d'un redémarrage prématuré. Des seuils mal calibrés — trop agressifs ou trop laxistes — génèrent soit des redémarrages en boucle, soit une détection de panne trop tardive.

Trois signaux, pas un seul

Une observabilité fiable combine systématiquement métriques (taux d'erreur, latence, saturation), logs centralisés et traces distribuées. Se reposer sur un seul de ces trois signaux laisse toujours un angle mort : les métriques signalent qu'un problème existe, les logs et les traces permettent de comprendre où et pourquoi.

COMPOSANT	RÔLE
Métriques (Prometheus)	Détection de dérive en temps réel — CPU, mémoire, latence, taux d'erreur par service
Logs centralisés	Reconstitution du contexte exact d'un incident, recherche transversale entre pods et nœuds
Traces distribuées	Localisation précise du service en cause dans une chaîne d'appels inter-services
Alerting	Notification automatique dès qu'un seuil critique est franchi, sans attendre une remontée manuelle

À retenir : l'objectif n'est pas de tout observer, mais de détecter en quelques minutes ce qui, sans instrumentation, ne serait découvert qu'au moment où un utilisateur signale un problème.

SECTION 06

Ressources, stockage et autoscaling

Un cluster sans limites de ressources définies fonctionne bien jusqu'au jour où un seul pod défaillant consomme la capacité de tous les autres.

Requêtes et limites, sur chaque conteneur

Chaque conteneur devrait déclarer des *requests* (ce qui lui est garanti) et des *limits* (ce qu'il ne peut pas dépasser) pour le CPU et la mémoire. Sans ces déclarations, l'ordonnanceur ne peut ni placer les pods intelligemment, ni protéger le cluster d'un service qui consommerait toute la capacité disponible sur un nœud.

Le trio d'autoscaling

Trois mécanismes d'autoscaling se combinent, chacun avec un rôle distinct : le **Horizontal Pod Autoscaler** ajuste le nombre de répliques selon la charge, le **Vertical Pod Autoscaler** ajuste les requêtes/limites d'un pod dans le temps, et le **Cluster Autoscaler** ajuste le nombre de nœuds disponibles. Les déployer isolément produit rarement le résultat attendu — c'est leur combinaison cohérente qui rend le cluster réellement élastique.

Stockage persistant et cycle de vie des données

Pour les charges avec état, les `PersistentVolumeClaim` doivent être associés à une classe de stockage adaptée à la charge réelle (IOPS, latence), avec une politique de sauvegarde testée — et pas seulement configurée. Un plan de sauvegarde jamais restauré en test n'est pas un plan de reprise fiable.

Piège fréquent : définir des limites de mémoire trop basses provoque des redémarrages en boucle (OOMKilled) difficiles à diagnostiquer si les logs ne conservent pas la trace du dépassement. Toujours dimensionner les limites à partir de mesures réelles, pas d'une estimation initiale.

SECTION 07

Grille d'auto-évaluation « prêt pour la prod »

Avant d'ouvrir un service à du trafic réel, vérifiez ces huit contrôles fondamentaux.

- ☐ **Chaque service critique tourne-t-il avec au moins 3 réplicas répartis sur plusieurs zones ?** Un réplica unique signifie une interruption garantie à la première panne de nœud.
- ☐ **Un PodDisruptionBudget protège-t-il chaque charge critique lors des mises à jour de nœud ?** Sans lui, une opération de maintenance planifiée peut faire tomber un service entier.
- ☐ **Les sondes liveness et readiness sont-elles calibrées sur des mesures réelles, pas des valeurs par défaut ?** Des seuils copiés-collés génèrent des redémarrages inutiles ou une détection trop tardive.
- ☐ **Le RBAC applique-t-il le moindre privilège, sans ClusterRoleBinding large ?** Un compte de service trop permissif transforme un pod compromis en accès cluster complet.
- ☐ **Des NetworkPolicy limitent-elles le trafic est-ouest entre espaces de noms ?** Sans elles, un seul pod compromis peut atteindre librement le reste du cluster.
- ☐ **Chaque conteneur déclare-t-il des requests et limits CPU/mémoire ?** Sans elles, un pod défaillant peut épuiser la capacité de tout un nœud.
- ☐ **Les métriques, logs et traces sont-ils centralisés et alertants ?** Sans ce trio, un incident n'est découvert qu'au moment où un utilisateur le signale.
- ☐ **Le plan de sauvegarde des volumes persistants a-t-il été restauré en test au moins une fois ?** Une sauvegarde jamais testée n'est pas une garantie de reprise.

Comment lire vos réponses : trois « non » ou plus sur ces huit contrôles indiquent qu'une revue de mise en production structurée devrait précéder toute ouverture à du trafic critique.

SECTION 08

Feuille de route en 5 étapes

Une démarche progressive, conçue pour sécuriser un cluster existant sans interrompre les services qu'il héberge déjà.

- 1 Auditer**
Passer chaque service critique au crible de la grille d'auto-évaluation en page 9, et documenter précisément les écarts identifiés.
- 2 Sécuriser en priorité**
Corriger d'abord le RBAC trop permissif et l'absence de politiques réseau — ce sont les écarts qui exposent le plus largement le cluster en cas de compromission d'un pod.
- 3 Fiabiliser la disponibilité**
Ajouter les réplicas manquants, définir les PodDisruptionBudget, et tester une coupure de nœud volontaire en pré-production.
- 4 Instrumenter**
Déployer le trio métriques/logs/traces sur les services qui en sont encore dépourvus, avec des alertes actionnables plutôt que des tableaux de bord passifs.
- 5 Automatiser et documenter**
Mettre en place l'autoscaling combiné (HPA/VPA/Cluster Autoscaler) et formaliser la checklist de mise en production pour tout nouveau service à venir.

CONCLUSION

Un cluster prêt pour la production, pas seulement pour la démo

La différence entre un cluster qui fonctionne et un cluster prêt pour la production ne se voit pas dans un déploiement réussi — elle se voit dans la façon dont le cluster absorbe une panne de nœud, une tentative d'intrusion ou un pic de trafic imprévu.

À mesure que l'adoption de Kubernetes progresse et que les charges d'IA s'y ajoutent, les huit contrôles détaillés dans ce document restent le socle sur lequel construire une plateforme fiable, quelle que soit la maturité de départ de l'équipe.

Move2Cloud	est un cabinet de conseil spécialisé en infrastructure, cloud et transformation IT. Nous accompagnons les organisations dans l'audit, la sécurisation et la mise en production de leurs clusters Kubernetes.
Contact	Move2Cloud — Paris · move2cloud.fr
Prochaine étape	Un audit de mise en production permet de passer chaque service critique au crible des huit contrôles de la page 9, avec un plan de remédiation priorisé.

Sources

CNCF, « Kubernetes Established as the De Facto Operating System for AI as Production Use Hits 82% in 2025 » — CNCF Annual Cloud Native Survey, édition 2026 — IBM X-Force Threat Intelligence Index (synthèse via deepstrike.io, « Cloud Security Statistics 2026 ») — Wiz, « Cloud Data Security Snapshot 2025 » — documentation Kubernetes officielle (RBAC, NetworkPolicy, PodDisruptionBudget, autoscaling).

Document rédigé à titre d'exemple et d'illustration par Move2Cloud à des fins pédagogiques. Les recommandations présentées sont générales ; leur mise en œuvre doit être adaptée au contexte technique et réglementaire de chaque organisation.