

Configuration d'un système Oracle® Solaris 11.2 en tant que routeur ou équilibreur de charge



Référence: E53804-02
Septembre 2014

Copyright © 2011, 2014, Oracle et/ou ses affiliés. Tous droits réservés.

Ce logiciel et la documentation qui l'accompagne sont fournis sous concédés sous licence et soumis à des restrictions d'utilisation et de divulgation et, sont protégés par les lois sur la propriété intellectuelle. Sauf disposition expresse de votre contrat de licence ou de la loi, vous ne pouvez pas copier, reproduire, traduire, diffuser, modifier, accorder de licence, transmettre, distribuer, exposer, exécuter, publier ou afficher le logiciel, même partiellement, sous quelque forme et par quelque procédé que ce soit. Par ailleurs, il est interdit de procéder à toute ingénierie inverse du logiciel, de le désassembler ou de le décompiler, excepté à des fins d'interopérabilité avec des logiciels tiers ou tel que prescrit par la loi.

Les informations fournies dans ce document sont susceptibles de modification sans préavis. Par ailleurs, Oracle Corporation ne garantit pas qu'elles soient exemptes d'erreurs et vous invite, le cas échéant, à lui en faire part par écrit.

Si ce logiciel, ou la documentation qui l'accompagne, est livré sous licence au Gouvernement des Etats-Unis, ou à quiconque qui aurait souscrit la licence de ce logiciel ou l'utilise pour le compte du Gouvernement des Etats-Unis, la notice suivante s'applique :

U.S. GOVERNMENT END USERS. Oracle programs, including any operating system, integrated software, any programs installed on the hardware, and/or documentation, delivered to U.S. Government end users are "commercial computer software" pursuant to the applicable Federal Acquisition Regulation and agency-specific supplemental regulations. As such, use, duplication, disclosure, modification, and adaptation of the programs, including any operating system, integrated software, any programs installed on the hardware, and/or documentation, shall be subject to license terms and license restrictions applicable to the programs. No other rights are granted to the U.S. Government.

Ce logiciel ou matériel a été développé pour un usage général dans le cadre d'applications de gestion des informations. Ce logiciel ou matériel n'est pas conçu ni n'est destiné à être utilisé dans des applications à risque, notamment dans des applications pouvant causer des dommages corporels. Si vous utilisez ce logiciel ou matériel dans le cadre d'applications dangereuses, il est de votre responsabilité de prendre toutes les mesures de secours, de sauvegarde, de redondance et autres mesures nécessaires à son utilisation dans des conditions optimales de sécurité. Oracle Corporation et ses affiliés déclinent toute responsabilité quant aux dommages causés par l'utilisation de ce logiciel ou matériel pour des applications dangereuses.

Oracle et Java sont des marques déposées d'Oracle Corporation et/ou de ses affiliés. Tout autre nom mentionné peut correspondre à des marques appartenant à d'autres propriétaires qu'Oracle.

Intel et Intel Xeon sont des marques ou des marques déposées d'Intel Corporation. Toutes les marques SPARC sont utilisées sous licence et sont des marques ou des marques déposées de SPARC International, Inc. AMD, Opteron, le logo AMD et le logo AMD Opteron sont des marques ou des marques déposées d'Advanced Micro Devices. UNIX est une marque déposée de The Open Group.

Ce logiciel ou matériel et la documentation qui l'accompagne peuvent fournir des informations ou des liens donnant accès à des contenus, des produits et des services émanant de tiers. Oracle Corporation et ses affiliés déclinent toute responsabilité ou garantie expresse quant aux contenus, produits ou services émanant de tiers. En aucun cas, Oracle Corporation et ses affiliés ne sauraient être tenus pour responsables des pertes subies, des coûts occasionnés ou des dommages causés par l'accès à des contenus, produits ou services tiers, ou à leur utilisation.

Table des matières

Utilisation de la présente documentation	7
 1 Présentation des routeurs et des équilibreurs de charge	9
Présentation du routeur	9
Protocoles de routage	10
Présentation du routeur VRRP	12
Présentation de l'équilibreur de charge intégré	13
Fonctions d'ILB	13
Pourquoi utiliser des routeurs et des équilibreurs de charge VRRP ?	15
 2 Configuration d'un système en tant que routeur	17
Configuration d'un routeur IPv4	17
▼ Configuration d'un routeur IPv4	17
Configuration d'un routeur IPv6	22
Démon in.ripngd, pour routage IPv6	22
Publications, préfixes et messages de routeur	22
▼ Configuration d'un routeur compatible IPv6	23
 3 Utilisation de Virtual Router Redundancy Protocol	27
A propos de VRRP	27
Fonctionnement du protocole VRRP	28
A propos de la fonction VRRP de couche 3	30
Comparaison de VRRP de couches 2 et 3	31
Limitations de routeur VRRP de couche 2 et 3	32
 4 Configuration et administration du protocole de redondance de routeur virtuel	35
Planification d'une configuration VRRP	35
Installation de VRRP	36
▼ Installation de VRRP	36

Configuration VRRP	36
Création d'une VNIC VRRP pour le VRRP de couche 2	37
Création d'un routeur VRRP	38
Configuration de l'adresse IP virtuelle pour les routeurs VRRP de couches 2 et 3	40
Activation et désactivation d'un routeur VRRP	41
Modification d'un routeur VRRP	42
Affichage des configurations de routeur VRRP de couches 2 et 3	42
Affichage des adresses IP associées à des routeurs VRRP	44
Suppression d'un routeur VRRP	45
Contrôle des messages Gratuitous ARP et NDP	45
Cas d'utilisation : Configuration d'un routeur VRRP de couche 2	46
5 Présentation d'un équilibreur de charge intégré	49
Composants ILB	49
Modes de fonctionnement d'ILB	50
Mode Direct Server Return (DSR)	50
Mode NAT (Network Address Translator)	51
Fonctionnement de l'équilibreur de charge intégré	54
6 Configuration et gestion de l'équilibreur de charge intégré	57
Installation d'ILB	57
Configuration d'ILB à l'aide de la CLI	58
Activation ou désactivation d'ILB	59
▼ Activation de l'équilibreur de charge intégré	59
▼ Désactivation de l'équilibreur de charge intégré	60
Gestion d'un ILB	60
Définition des groupes de serveurs et des serveurs backend dans ILB	61
Surveillance de l'état de fonctionnement dans ILB	64
Configuration des règles ILB	67
Cas d'utilisation : configuration d'un ILB	70
Affichage des statistiques ILB	71
Affichage des données statistiques	72
Affichage de la table de connexions NAT	72
Affichage de la table de mise en correspondance de la persistance de session	73
Importation et exportation des configurations	74
7 Configuration ILB pour une haute disponibilité	75
Configuration d'ILB pour la haute disponibilité à l'aide de la topologie DSR	75

▼ Configuration d'ILB pour la haute disponibilité à l'aide de la topologie DSR	77
Configuration d'ILB pour la haute disponibilité à l'aide de la topologie Half-NAT	78
▼ Configuration d'ILB pour la haute disponibilité à l'aide de la topologie Half-NAT	80
Index	83

Utilisation de la présente documentation

- **Présentation:** Décrit la procédure de configuration en tant qu'Oracle Solaris 11.2 IPv4 ou d'un routeur IPv6. Fournit une vue d'ensemble et des instructions de configuration de Protocole de redondance de routeur virtuel (VRRP) et de l'équilibreur de charge intégré (ILB).
- **Public visé :** administrateurs système.
- **Connaissances requises :** connaissances en matière d'administration réseau de base et avancées.

Bibliothèque de la documentation des produits

Les informations de dernière minute et les problèmes connus pour ce produit sont inclus dans la bibliothèque de documentation accessible à l'adresse <http://www.oracle.com/pls/topic/lookup?ctx=E36784>.

Accès aux services de support Oracle

Les clients Oracle ont accès au support électronique via My Oracle Support. Pour plus d'informations, visitez le site <http://www.oracle.com/pls/topic/lookup?ctx=acc&id=info> ou le site <http://www.oracle.com/pls/topic/lookup?ctx=acc&id=trs> si vous êtes malentendant.

Commentaires en retour

Faites part de vos commentaires sur cette documentation à l'adresse : <http://www.oracle.com/goto/docfeedback>.

Présentation des routeurs et des équilibreurs de charge

Ce chapitre décrit la méthode utilisée par les routeurs sont utilisés dans et programmes d'équilibrage de charge pour la connexion réseau d'ordinateurs s'Oracle Solaris et de distribuer charges de travail. Les routeurs gèrent les activités de routage à l'aide de protocoles tels que RIP (Routing Information Protocol), RIPng (next generation IP), RDISC (Internet Control Message Protocol Router Discovery), OSPF (Open Shortest Path First), BGP (Border Gateway Protocol), IS-IS (Intermediate System to Intermediate System) et VRRP (Virtual Router Redundancy Protocol).

Programmes d'équilibrage de charge distribuer le trafic réseau sur un certain nombre de serveurs. La ventilation d'un réseau optimale charge de travail du contribue à atteindre pour accroître le débit de partage de ressources et la disponibilité et.

Ce chapitre aborde les sujets suivants :

- [“Présentation du routeur” à la page 9](#)
- [“Présentation du routeur VRRP” à la page 12](#)
- [“Présentation de l'équilibreur de charge intégré” à la page 13](#)
- [“Pourquoi utiliser des routeurs et des équilibreurs de charge VRRP ?” à la page 15](#)

Présentation du routeur

Un routeur correspond à un périphérique utilisé dans un réseau informatique pour connecter des ordinateurs et de transfert de paquets de données entre les ordinateurs du réseau. Un routeur peuvent avoir deux ou plusieurs connexions à partir de différents réseaux. Lit les informations d'adresse, le routeur from the incoming afin de déterminer leur de données de destination des paquets de données. Ensuite, les paquets sont envoyés à l'autre réseau à l'aide des informations contenues dans la table de routage du routeur. Processus selon lequel les routeurs ce trafic faisant pointer est reproduit jusqu'à ce que le noeud ceux-ci atteignent la destination des données.

Protocoles de routage

Les protocoles de routage servent à traiter l'activité de routage dans un système. Les routeurs échangent des informations avec les autres hôtes afin de préserver les routes connues vers les réseaux distants. Les routeurs et les hôtes peuvent exécuter des protocoles de routage. Les protocoles de routage sur l'hôte communiquent avec les démons de routage sur d'autres routeurs et hôtes. Ces protocoles aident l'hôte à identifier la destination des paquets. Lorsque les interfaces réseau sont activées, le système communique automatiquement avec les démons de routage. Ceux-ci surveillent les routeurs sur le réseau et signalent les adresses des routeurs aux hôtes sur le réseau local. Certains protocoles de routage conservent également des statistiques permettant de mesurer les performances du routage. Comme pour la transmission de paquets, vous devez configurer de manière explicite le routage sur un système Oracle Solaris.

RIP et RDISC constituent des protocoles TCP/IP standard. Le tableau suivant décrit les protocoles de routage dans Oracle Solaris pris en charge.

TABEAU 1-1 Protocoles de routage Oracle Solaris

Protocole	Démon associé	Description	Pour obtenir des instructions
RIP	in.routed	IGP (Interior Gateway Protocol), qui route les paquets IPv4 et gère une table de routage	“Configuration d'un routeur IPv4” à la page 17
RDISC	in.routed	Permet aux hôtes de détecter la présence d'un routeur sur le réseau	“ Configuration du routage de systèmes à interface unique ” du manuel “ Configuration et administration des composants réseau dans Oracle Solaris 11.2 ”
RIPng	in.ripngd	IGP acheminant les paquets IPv6 et gérant une table de routage	“Configuration d'un routeur compatible IPv6” à la page 23
Neighbor Discovery Protocol (NDP)	in.ndpd	Signale la présence d'un routeur IPv6 et détecte les hôtes IPv6 sur un réseau	“ Configuration d'un système pour IPv6 ” du manuel “ Configuration et administration des composants réseau dans Oracle Solaris 11.2 ”

Pour plus d'informations sur les tables et types de routage, reportez-vous à la section [“ Tables et types de routage ” du manuel “ Configuration et administration des composants réseau dans Oracle Solaris 11.2 ”](#).

RIP (Routing Information Protocol)

Le protocole d'information de routage (RIP) est un protocole de routage sur un vecteur de distance. Utilise un compteur RIP parle de gamme de mesures du prochain saut. Il est implémenté par le démon de routage in.routed. Le démon démarre automatiquement lors de l'initialisation du système. En cas d'exécution sur un routeur avec l'option -s spécifiée, le démon in.routed renseigne la table de routage du noyau en indiquant une route pour chaque réseau

accessible et publie l'accessibilité via toutes les interfaces réseau. Lorsqu'il est exécuté sur un hôte avec l'option `-q` spécifiée, le démon `in.routed` extrait les informations de routage mais ne publie pas l'accessibilité.

Sur les hôtes, vous pouvez extraire les informations de routage de deux façons :

- Quand le drapeau *est pas* spécifié (`S` majuscule ou mode économie d'espace). Le démon `in.routed` construit une table de routage complète exactement de la même manière que sur un routeur.
- En spécifiant l'indicateur. Le démon `in.routed` crée une table de routage minimale pour le noyau, contenant une seule route par défaut pour chaque routeur disponible.

Protocole de détection de routeur ICMP

Les hôtes utilisent le protocole de détection de routeur (RDISC) pour obtenir les informations de routage de la part des routeurs. Par conséquent, lorsque les hôtes exécutent RDISC, les routeurs doivent également exécuter un autre protocole, par exemple RIP, afin d'échanger les informations de routeur.

Le démon RDISC est implémenté par `in.routed` qui doit s'exécuter à la fois sur les routeurs et sur les hôtes. Sur les hôtes, `in.routed` utilise RDISC pour détecter les routes par défaut des routeurs qui se publient eux-mêmes via RDISC. Sur les routeurs, `in.routed` utilise RDISC pour publier les routes par défaut des hôtes sur les réseaux directement connectés. Voir la page de manuel [in.routed\(1M\)](#) man page and the [gateways\(4\)](#) pour plus d'informations.

Suite de protocoles de routage Quagga

Suite Quagga est un logiciel de routage RIP qui, permet la mise en oeuvre, ainsi que le nombre d'ouvertures de RIPng OSPF (système) System, Intermédiaire-sur Intermédiaire, IS IS (Border Gateway Protocol) BGP et protocoles (y compris) Oracle Solaris pour les plates-formes UNIX.

RIPng offre une extension de RIP pour la prise en charge d'IPv6, y compris diverses améliorations d'IPv6. Les fonctions RIPng sont similaires à celles de RIP.

Protocole OSPF est un routeur gamme qui est utilisé pour distribuer Autonomous System, les informations réseau dans un plus grand. La dernière version de OSPF, est toujours pris en charge, IPv6 OSPFv3.

IS-IS est un protocole de routage dynamique d'état de lien qui est utilisé pour distribuer les informations relatives aux gammes de fournisseur de service réseau à l'intérieur d'une grande.

Utilise un ensemble de BGP réseaux préfixé pour rendre les décisions de routage IP basé sur celui d'un système autonome et les règles entre les réseaux de grande taille.

Le tableau suivant répertorie la suite de protocoles de routage Quagga Open Source qu'Oracle Solaris prend également en charge.

TABEAU 1-2 Suite de protocoles de routage Quagga

Protocole	Démon associé	Description
RIP	ripd	Protocole IGP à vecteur de distance IPv4 qui achemine les paquets IPv4 et signale sa table de routage aux routeurs adjacents
RIPng	ripngd	Distance vectoring IPv6 IGP acheminant les paquets IPv6 et gérant une table de routage
OSPF	ospfd	Protocole IGP d'état des liens IPv4 pour le routage des paquets et la mise en réseau à haute disponibilité.
BGP	bgpd	IPv4 et IPv6 (Extérieure EGP Gateway Protocol) pour le routage d'un domaine administratif à
IS-IS	isisd	L'état du lien IPv4 et IPv6 pour le routage IGP ou réseau au sein d'un domaine d'administration

Pour plus d'informations sur les protocoles Quagga, visitez le site Web de Quagga Routing Suite à l'adresse <http://www.nongnu.org/quagga/index.html>.

Protocole de redondance de routeur virtuel

IP permet d'accroître la disponibilité de VRRP, tels que ceux qui les adresses des routeurs et programmes d'équilibrage de charge utilisé. VRRP est un protocole internet standard spécifié dans le document [Virtual Router Redundancy Protocol 5798, RFC Version 3 for IPv4 and IPv6](#). Fournit un outil d'administration Oracle Solaris qui configure et gère le service VRRP.

En plus de VRRP standard et existants, au niveau de la couche 2 une clause de Oracle Solaris 11.2 VRRP fournit la prise en charge Layer 3 sur IP (IPMP) et VRRP les interfaces InfiniBand VRRP et une prise en charge améliorée de VRRP dans les zones.

Pour plus d'informations sur l'utilisation des routeurs VRRP et la configuration, reportez-vous au [Chapitre 3, Utilisation de Virtual Router Redundancy Protocol](#) et au [Chapitre 4, Configuration et administration du protocole de redondance de routeur virtuel](#).

Présentation du routeur VRRP

Le document Virtual Router Redundancy Protocol routeur une image de routeur unique est créé par le fonctionnement d'un ou plusieurs routeurs qui utilisent le VRRP. Le service VRRP est exécuté sur chaque routeur VRRP et gère son état. Un hôte peut compter plusieurs routeurs VRRP configurés, chacun correspondant à un routeur virtuel différent.

Il utilise la couche 2 VRRP et exige la norme unique protocole adresse MAC de routeur virtuel MAC. Les adresses IP virtuel soit résolu vers le même virtuel toujours adresse MAC. Vous devez créer un MAC permet d'obtenir la VNIC VRRP adresse MAC de routeur virtuel unique. La fonctionnalité VRRP d 'propriétaire de couche 3 évite les Oracle Solaris pour configurer complètement unique VRRP adresse MAC virtuelle VRRP et, par conséquent, pour la prise en charge des routeurs VRRP IPMP fournit et pour les interfaces InfiniBand sur.

Pour plus d'informations sur l'utilisation des routeurs VRRP et la configuration, reportez-vous au [Chapitre 3, Utilisation de Virtual Router Redundancy Protocol](#) et au [Chapitre 4, Configuration et administration du protocole de redondance de routeur virtuel](#).

Présentation de l'équilibreur de charge intégré

Dans Oracle Solaris l'équilibreur de charge intégré (ILB, IP Network Multipathing) permet de charger des couches 3 et l'équilibrage de charge de type Layer 4. ILB fonctionne au niveau des couches réseau (IP) et de transport (TCP/UDP) pour le système d'exploitation Oracle Solaris installé sur des systèmes SPARC et x86. Peut être utilisée pour améliorer ILB de fiabilité et d'évolutivité, et pour réduire le temps de réponse des services réseau.

ILB intercepte les demandes entrantes des clients, détermine le serveur d'arrière-plan qui doit gérer la demande en fonction des règles d'équilibrage de charge, puis transfère la demande au serveur sélectionné. ILB peut également être utilisé en tant que routeur pour le serveur d'arrière-plan . ILB effectue éventuellement des vérifications d'état et fournit les données de sorte que les algorithmes d'équilibrage de la charge vérifient si le serveur sélectionné peut traiter la demande entrante.

Fonctions d'ILB

ILB possède les fonctions principales suivantes :

- Prise en charge des modes DSR (Direct Server Return) et NAT (Network Address Translation) sans conservation de statut pour IPv4 et IPv6.
Pour plus d'informations sur les modes de fonctionnement DSR et NAT, reportez-vous à la section [“Modes de fonctionnement d'ILB” à la page 50](#).
- Assiste le trafic et répartir la charge et sélectionner les serveurs à l'aide d'un ensemble d'algorithmes pour chaque mode de fonctionnement.
- Administration de l'équilibreur de charge intégré dans une interface de ligne de commande (CLI)

Pour plus d'informations sur la configuration ILB, reportez-vous à [“Configuration d'ILB à l'aide de la CLI” à la page 58](#).

- Capacités de surveillance des serveurs par le biais de vérifications de l'état.

Pour plus d'informations sur la capacité de surveillance du serveur, reportez-vous à [“Surveillance de l'état de fonctionnement dans ILB” à la page 64.](#)

Le tableau suivant répertorie et décrit les fonctionnalités d'ILB qui sont disponibles pour les différents modes de fonctionnement.

TABEAU 1-3 Fonctionnalités ILB

Fonctionnalités de	Description	Modes de fonctionnement
Envoi de la commande ping aux adresses IP virtuelles (VIP)	Répond aux demandes d'écho ILB provenant des clients VIP ICMP adresses IP virtuelles.	Les modes DSR et NAT à la fois
Vous permet d'ajouter ou enlever des serveurs d'un groupe sans interruption de service	Permet d'enlever des serveurs ILB ajoute de façon dynamique du groupe de serveurs ou.	Mode NAT
Permet de configurer la persistance de session (" caractère résident / non permanent ")	ILB permet de configurer la persistance de session pour les applications ou pour envoyer les paquets d'un client les connexions au même serveur d'arrière-plan . Vous pouvez configurer la persistance de session (c'est-à-dire la conservation de l'adresse source) pour un service virtuel en ajoutant l'option -p et en spécifiant le masque de réseau pmask dans la commande <code>ilbadm create-rule</code> . Pour plus d'informations, reportez-vous à “Création d'une règle de ILB” à la page 68.	Les modes DSR et NAT à la fois
Un drainage automatique vous permet d'effectuer la connexion	Cette capacité empêche l'établissement de nouvelles connexions à un serveur désactivé. Cette fonction est utile car elle arrête les serveurs sans interrompre les connexions ou sessions actives. Les connexions existantes continuent de fonctionner. Après avoir mis fin à toutes ses connexions, vous pouvez arrêter un serveur en vue d'effectuer des opérations de maintenance. Dès que le serveur est à nouveau en mesure de traiter des demandes, il faut l'activer pour que l'équilibreur de charge établisse de nouvelles connexions avec lui.	Mode NAT
Equilibrage de la charge des TCP (Transmission Control Protocol) et UDP (User Datagram Protocol)	ILB peut équilibrer la charge de tous les ports sur une adresse IP donnée entre différents ensembles de serveurs sans nécessiter la configuration de règles explicites pour chaque port.	Les modes DSR et NAT à la fois
Permet de spécifier chaque port pour les services virtuels au sein du même groupe de serveurs	Vous permet de spécifier différentes ILB des ports de destination à des serveurs appartenant au même groupe de serveurs.	Mode NAT
Cette fonctionnalité vous permet de charger-une plage de ports simple	Charge ILB sur une fourchette de soldes par le biais de ports VIP pour un groupe de serveurs donné. Il est pratique de conserver les adresses IP en équilibrant la charge des différentes plages de ports associées à la même adresse VIP sur divers ensembles de serveurs backend. De plus, lorsque la persistance de session est activée en mode NAT, ILB envoie au même serveur backend les demandes provenant de la même adresse IP client à destination de différents ports de la plage.	Les deux modes DSR une NAT
Permet de changer de port et	Plage de ports ces fonction dépendent de la plage de ports d'un serveur dans une règle d'équilibrage de charge. Si la	Mode NAT

Fonctionnalités de	Description	Modes de fonctionnement
	plage de ports d'un serveur est différente de la plage de ports VIP, la fonction de basculement entre les ports est automatiquement implémentée. Quant à la fonction de réduction de la plage de ports, elle est implémentée si la plage de ports du serveur se limite à un seul port.	

Pour plus d'informations sur les composants ILB, les modes de fonctionnement I, les algorithmes et la façon dont fonctionne ILB, reportez-vous au [Chapitre 5, Présentation d'un équilibreur de charge intégré](#). Pour plus d'informations sur la configuration et la gestion ILB, reportez-vous au [Chapitre 6, Configuration et gestion de l'équilibreur de charge intégré](#) et [Chapitre 7, Configuration ILB pour une haute disponibilité](#).

Pourquoi utiliser des routeurs et des équilibreurs de charge VRRP ?

Lorsque vous configurez un réseau, comme un réseau local (LAN), il est très important de fournir un service de haute disponibilité. La haute disponibilité est un état rendant prend le relais en cas de redondant, reportez-vous à la système au moment de l'échec afin de garantir la continuité de l'activité. La haute disponibilité est également applicable lorsque trop charge globale est déportée sur redondant, reportez-vous à la système. L'état Peut faire l'objet d'une part importante du haute disponibilité dans les situations comme un arrêt prévu ou l'indisponibilité imprévue, l'équilibrage des charges, et la reprise sur sinistre.

Au sein d'un domaine réseau, il est possible de mettre en oeuvre de haute disponibilité à différents niveaux, par exemple, lien, IP et routeurs. Les routeurs de lecture programmes d'équilibrage de charge et un rôle important qu'elle offre un service de haute disponibilité. Les routeurs VRRP dans Oracle Solaris sont, au réseau et niveau ILB le basculement en cas d'incident et de partage de charge des mécanismes permettant d'assurer la haute disponibilité.

Pour plus d'informations sur l'utilisation et la configuration des routeurs VRRP, reportez-vous au [Chapitre 3, Utilisation de Virtual Router Redundancy Protocol](#). Pour plus d'informations sur l'utilisation et la configuration d'ILB, reportez-vous au [Chapitre 5, Présentation d'un équilibreur de charge intégré](#) et au [Chapitre 6, Configuration et gestion de l'équilibreur de charge intégré](#).

Configuration d'un système en tant que routeur

Un routeur fournit l'interface entre deux réseaux ou plus. Par conséquent, vous devez attribuer un nom unique et une adresse IP à chacune des interfaces de réseau physique du routeur. Par conséquent, chaque routeur possède un nom d'hôte et une adresse IP associés à son interface réseau principale ainsi qu'un nom et une adresse IP uniques pour chaque interface réseau supplémentaire. Ce chapitre explique comment configurer le système Oracle Solaris d'un routeur IPv4 ou sous la forme d'un routeur IPv6.

Ce chapitre aborde les sujets suivants :

- [“Configuration d'un routeur IPv4” à la page 17](#)
- [“Configuration d'un routeur IPv6” à la page 22](#)

Pour plus d'informations sur la configuration d'hôte gamme de fabrication à un réseau Oracle Solaris, reportez-vous à [“ Configuration du routage de systèmes à interface unique ” du manuel “ Configuration et administration des composants réseau dans Oracle Solaris 11.2 ”](#).

Pour obtenir des informations de référence sur les protocoles de routage, reportez-vous à [“Protocoles de routage” à la page 10](#) et [“ A propos du routage IPv6 ” du manuel “ Configuration et administration des composants réseau dans Oracle Solaris 11.2 ”](#).

Configuration d'un routeur IPv4

Vous pouvez également effectuer la procédure suivante pour configurer un système doté d'une seule interface physique (un hôte, par défaut) en tant que routeur. Vous pouvez configurer un système d'interface unique en tant que routeur si le système fait office de point d'extrémité sur un lien PPP, comme décrit à la section [“ Planification d’une liaison PPP commutée ” du manuel “ Gestion de réseaux série à l’aide d’UUCP et de PPP dans Oracle Solaris 11.2 ”](#).

▼ Configuration d'un routeur IPv4

La procédure suivante suppose que vous configurez les interfaces du routeur après l'installation.

Avant de commencer

Une fois le routeur physiquement installé sur le réseau, configurez le routeur de sorte qu'il fonctionne en mode fichiers locaux. Cette configuration garantit l'initialisation du routeur en cas de panne du serveur de configuration.

1. Connexion en tant qu'administrateur.

Pour plus d'informations, reportez-vous à la section [“ A l'aide de vos droits administratifs attribués ” du manuel “ Sécurisation des utilisateurs et des processus dans Oracle Solaris 11.2 ”](#).

2. Configurez les interfaces IP sur les cartes réseau du système.

```
# ipadm create-ip IP-interface
```

3. Configurez l'interface IP avec une adresse IP valide, vous devez choisir l'une des commandes suivantes.**■ Pour configurer une adresse statique, saisissez la commande suivante :**

```
# ipadm create-addr -a address [interface | addr-obj]
```

■ Pour configurer une adresse non statique, tapez la commande suivante :

```
# ipadm create-addr -T address-type [interface | addr-obj]
```

Pour obtenir des instructions détaillées sur la procédure de configuration, reportez-vous au [Chapitre 3, “ Configuration et administration des interfaces et adresses IP dans Oracle Solaris ”](#) du manuel [“ Configuration et administration des composants réseau dans Oracle Solaris 11.2 ”](#).

Assurez-vous que chaque interface IP est configurée avec l'adresse IP du réseau pour lequel le système acheminera les paquets. Par conséquent, si le système sert les réseaux 192.168.5.0 et 10.0.5.0, il faut configurer une NIC pour chaque réseau.



Attention - Assurez - vous de l'administration DHCP de manière approfondie avant de configurer en charge un routeur IPv4 de manière à utiliser DHCP.

4. Ajoutez le nom d'hôte et l'adresse IP de chaque interface au fichier `/etc/inet/hosts`.

Par exemple, supposons que les noms assignés aux deux interfaces du routeur sont respectivement `krakatoa` et `krakatoa-1`. Les entrées dans le fichier `/etc/inet/hosts` seraient comme suit :

```
192.168.5.1      krakatoa      #interface for network 192.168.5.0
10.0.5.1        krakatoa-1    #interface for network 10.0.5.0
```

5. Effectuez la procédure [“ Configuration d'un système en mode fichiers locaux ”](#) du manuel [“ Configuration et administration des composants réseau dans Oracle Solaris 11.2 ”](#).

6. **Si le routeur est connecté à un réseau comportant des sous-réseaux, ajoutez le numéro de réseau et le masque de réseau au fichier `/etc/inet/netmasks`.**

Par exemple, pour adresse IPv4 de numérotation, telle que 192.168.5.0, tapez ce qui suit :

```
192.168.5.0    255.255.255.0
```

7. **Activez le transfert de paquets IPv4 sur le routeur.**

```
# ipadm set-prop -p forwarding=on ipv4
```

8. **(Facultatif) Lancez un protocole de routage.**

Utilisez l'une des commandes suivantes :

```
# routeadm -e ipv4-routing -u
```

où l'option `-e` active le routage IPv4 où et l'option `option-u` s'applique à jour la configuration en cours dans le système en cours d'exécution.

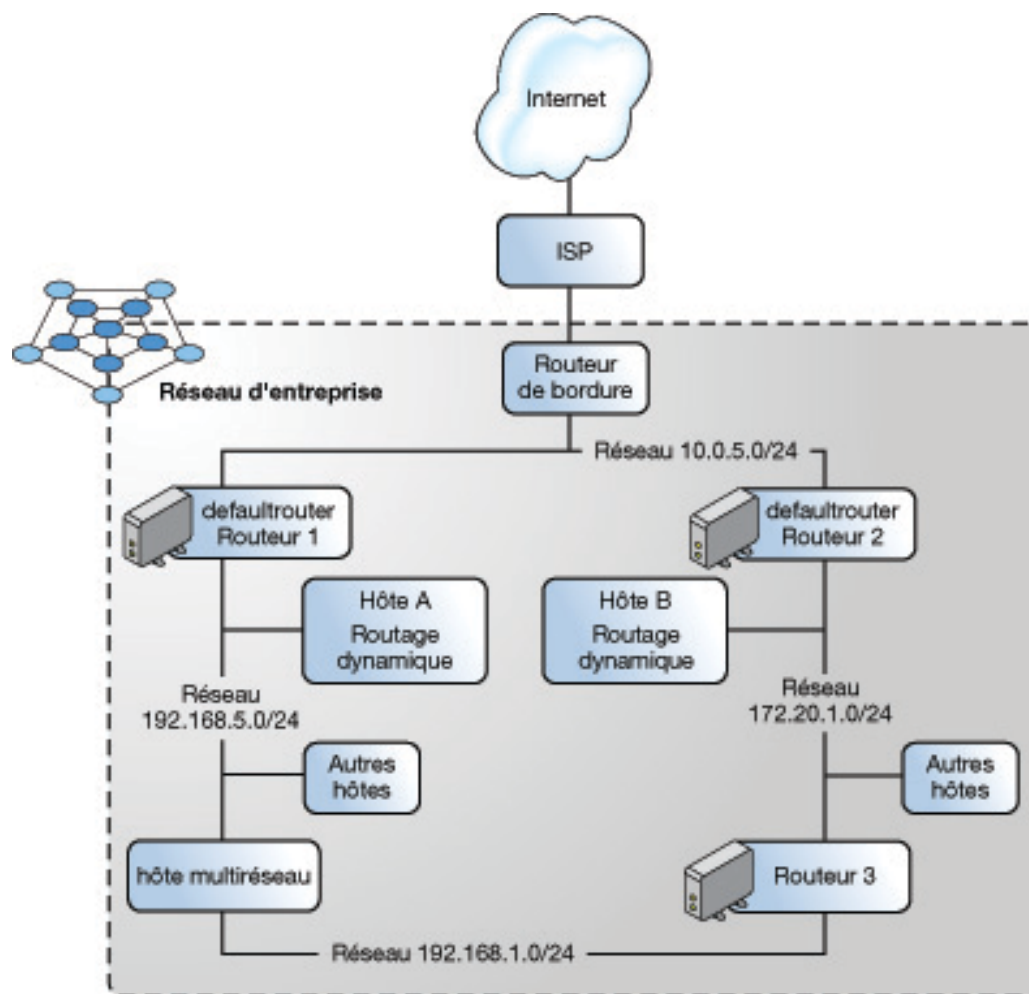
```
# svcadm enable route:default
```

Lorsque vous démarrez un protocole de routage, le démon de routage `/usr/sbin/in.routed` met automatiquement à jour la table de routage ; il s'agit d'un processus appelé *routage dynamique*. Pour plus d'informations sur les types de routage, reportez-vous à “ [Tables et types de routage](#) ” du manuel “ [Configuration et administration des composants réseau dans Oracle Solaris 11.2](#) ”. Pour plus d'informations sur la commande `routeadm`, reportez-vous à la page de manuel [routeadm\(1M\)](#), et pour obtenir plus d'informations sur la commande `ipadm`, reportez-vous à la page de manuel [ipadm\(1M\)](#).

L'utilitaire de gestion des services (SMF) FMRI (Fault Management Resource Identifiers) associé au démon `in.routed` est `svc:/network/routing/route`.

Exemple 2-1 Configuration d'un système en tant que routeur

Cet exemple de configuration repose sur les hypothèses suivantes :



Le routeur 2 contient deux connexions réseau câblées, une connexion au réseau 172.20.1.0 et une connexion au réseau 10.0.5.0. L'exemple montre comment configurer et installer un système en tant que routeur (routeur 2) du réseau 172.20.1.0. L'exemple suppose également que le routeur 2 a été configuré pour fonctionner dans le mode de fichiers locaux comme décrit dans la section [“ Configuration d'un système en mode fichiers locaux ”](#) du manuel [“ Configuration et administration des composants réseau dans Oracle Solaris 11.2 ”](#).

1. Déterminer le statut des interfaces du système.

```
# dladm show-link
```

LINK	CLASS	MTU	STATE	BRIDGE	OVER
net0	phys	1500	up	--	--
net1	phys	1500	up	--	--
net2	phys	1500	up	--	--

```
# ipadm show-addr
ADDROBJ      TYPE      STATE      ADDR
lo0/v4        static    ok          10.0.0.1/8
net0/v4        static    ok          172.20.1.10/24
```

2. Seule net0 a été configurée avec une adresse IP. Pour que le routeur 2 soit le routeur par défaut, vous devez connecter l'interface net1 physiquement au réseau 10.0.5.0.

```
# ipadm create-ip net1
# ipadm create-addr -a 10.0.5.10/24 net1
# ipadm show-addr
ADDROBJ      TYPE      STATE      ADDR
lo0/v4        static    ok          192.168.0.1/8
net0/v4        static    ok          172.20.1.10/24
net1/v4        static    ok          10.0.5.10/24
```

3. Ensuite, mettez à jour les bases de données réseau suivantes à l'aide des informations sur l'interface que vous venez de configurer et le réseau auquel elle est connectée.

```
# pfedit /etc/inet/hosts
192.168.0.1      localhost
172.20.1.10     router2        #interface for network 172.20.1
10.0.5.10       router2-out    #interface for network 10.0.5
# pfedit /etc/inet/netmasks
172.20.1.0      255.255.255.0
10.0.5.0        255.255.255.0
```

4. Enfin, activez le transfert de paquets ainsi que le démon de routage in.routed.

```
# ipadm set-prop -p forwarding=on ipv4
# svcadm enable route:default
```

Le transfert de paquets IPv4 et le routage dynamique via RIP sont maintenant activés sur le routeur 2. Cependant, pour terminer la configuration du routeur par défaut pour le réseau 172.20.1.0, vous devez effectuer les opérations suivantes :

- Modifiez les hôtes du réseau 172.20.1.0 pour qu'ils reçoivent leurs informations de routage du nouveau routeur par défaut. Pour plus d'informations, reportez-vous à [“ Création de routes permanentes \(statiques\) ”](#) du manuel [“ Configuration et administration des composants réseau dans Oracle Solaris 11.2 ”](#).
- Définissez une route statique menant au routeur de bordure dans la table de routage du Routeur 2. Pour plus de détails, reportez-vous à [“ Tables et types de routage ”](#) du manuel [“ Configuration et administration des composants réseau dans Oracle Solaris 11.2 ”](#). Pour plus d'informations sur la commande ipadm, reportez-vous à la page de manuel [ipadm\(1M\)](#).

Configuration d'un routeur IPv6

Cette section décrit les tâches de configuration d'un routeur IPv6.

Démon `in.ripngd`, pour routage IPv6

Le `in.ripngd` démon implémente le protocole d'information de routage de la prochaine génération pour routeurs IPv6 (RIPng). Le RIPng définit l'équivalent IPv6 de RIP (Routing Information Protocol, protocole d'informations de routage). Lorsque vous configurez un routeur IPv6 avec la commande `routeadm` et activez le routage IPv6, le démon `in.ripngd` implémente RIPng sur le routeur. Pour plus d'informations sur les options RIPng prises en charge, reportez-vous à [in.ripngd\(1M\)](#).

Publications, préfixes et messages de routeur

Sur des liens compatibles multicast et des liens point à point, chaque routeur envoie régulièrement un paquet de publication de routeur au groupe multicast pour lui annoncer sa disponibilité. Un hôte reçoit des publications de routeur de la totalité des routeurs, constituant une liste des routeurs par défaut. Les routeurs génèrent des publications de routeur de façon suffisamment fréquente pour permettre aux hôtes d'être avertis de leur présence en quelques minutes. Cependant, les routeurs n'effectuent pas de publications à une fréquence suffisante pour se fier à une absence de publication permettant de détecter une défaillance de routeur. Un algorithme de détection séparé qui détermine l'inaccessibilité de voisin fournit la détection de défaillance.

Les publications de routeur contiennent une liste de préfixes de sous-réseau utilisés pour déterminer si un hôte se trouve sur le même lien que le routeur. La liste de préfixes est également utilisée pour la configuration d'adresses autonomes. Les indicateurs associés aux préfixes spécifient les utilisations spécifiques d'un préfixe particulier. Les hôtes utilisent les préfixes sur liaison publiés afin de constituer et de maintenir une liste utilisée pour décider lorsque la destination d'un paquet se trouve sur la liaison ou au-delà d'un routeur. Une destination peut se trouver sur une liaison même si celle-ci n'est couverte par aucun préfixe sur liaison publié. Dans de tels cas, un routeur peut envoyer une redirection. La redirection informe l'expéditeur que la destination est un voisin.

Les publications de routeur et les indicateurs par préfixe permettent aux routeurs d'informer des hôtes de la méthode qu'ils doivent utiliser pour effectuer une configuration automatique d'adresse sans état.

Les messages de publication de routeur contiennent également des paramètres Internet, comme la limite de saut, que les hôtes devraient utiliser dans des paquets entrants. Les messages de

publication de routeur peuvent également (facultativement) contenir des paramètres de liens, comme le lien MTU. Cette fonctionnalité permet l'administration centralisée des paramètres critiques. Les paramètres peuvent être définis sur des routeurs et propagés automatiquement à tous les hôtes qui y sont connectés.

Les noeuds effectuent la résolution d'adresses par l'envoi de sollicitation de voisin à un groupe multicast, demandant au noeud cible de retourner son adresse de couche liaison. Les messages de sollicitation de voisin multicast sont envoyés à l'adresse de noeud multicast demandée de l'adresse cible. La cible retourne son adresse de couche liaison dans un message de publication d'un voisin unicast. Une paire de paquets de demande-réponse unique est suffisante pour permettre à l'initiateur et à la cible de résoudre les adresses de couche liaison de l'un et de l'autre. L'initiateur inclut son adresse de couche liaison dans la sollicitation de voisin.

▼ Configuration d'un routeur compatible IPv6

La procédure suivante présuppose que vous avez déjà configuré le système pour IPv6. Pour plus d'informations sur les procédures, reportez-vous au [Chapitre 3, “ Configuration et administration des interfaces et adresses IP dans Oracle Solaris ”](#) du manuel “ Configuration et administration des composants réseau dans Oracle Solaris 11.2 ”.

1. Connexion en tant qu'administrateur.

Pour plus d'informations, reportez-vous à la section “ [A l'aide de vos droits administratifs attribués](#) ” du manuel “ [Sécurisation des utilisateurs et des processus dans Oracle Solaris 11.2](#) ”.

2. Configuration d'un transfert de paquet IPv6 sur toutes les interfaces du routeur.

```
# ipadm set-prop -p forwarding=on ipv6
```

3. Démarrez le démon de routage.

Le démon `in.ripngd` gère le routage IPv6. Pour activer le protocole IPv6 d'un acheminement par exécuter l'une des commandes suivantes :

- Utilisez la commande `routeadm` :

```
# routeadm -e ipv6-routing -u
```

où l'option `-e` active le routage IPv4 où et l'option `option-u` s'applique à jour la configuration en cours dans le système en cours d'exécution.

- Utilisez la commande SMF adéquate :

```
# svcadm enable ripng:default
```

Pour plus d'informations sur la commande `routeadm`, reportez-vous à la page de manuel [routeadm\(1M\)](#).

4. Créez le fichier `/etc/inet/ndpd.conf`.

Spécifiez le préfixe de site que doit publier le routeur et les autres informations de configuration dans `/etc/inet/ndpd.conf`. Ce fichier est lu par le démon `in.ndpd`, qui implémente le protocole de détection de voisins IPv6.

Pour une liste des variables et valeurs autorisées, reportez-vous à la page du manuel [ndpd.conf\(4\)](#).

5. Saisissez le texte suivant dans le fichier `/etc/inet/ndpd.conf` :

```
ifdefault AdvSendAdvertisements true
prefixdefault AdvOnLinkFlag on AdvAutonomousFlag on
```

Ce texte indique au démon `in.ndpd` qu'il doit envoyer les publications de routeur à toutes les interfaces du routeur qui sont configurées pour IPv6.

6. Pour configurer le préfixe de site sur les différentes interfaces du routeur, ajoutez du texte supplémentaire au fichier `/etc/inet/ndpd.conf`.

Le texte doit être pris en compte dans le format suivant :

```
prefix global-routing-prefix:subnet ID/64 interface
```

Le fichier d'exemple `/etc/inet/ndpd.conf` suivant configure le routeur de sorte qu'il publie le préfixe de site `2001:0db8:3c4d::/48` sur les interfaces `net0` et `net1`.

```
ifdefault AdvSendAdvertisements true
prefixdefault AdvOnLinkFlag on AdvAutonomousFlag on
```

```
if net0 AdvSendAdvertisements 1
prefix 2001:0db8:3c4d:15::0/64 net0
```

```
if net1 AdvSendAdvertisements 1
prefix 2001:0db8:3c4d:16::0/64 net1
```

7. Réinitialisez le système.

Le routeur IPv6 commence la publication sur la liaison locale de tout préfixe de site dans le fichier `ndpd.conf`.

8. Afficher l'interface configurée pour IPv6.

```
# ipadm show-addr
ADDROBJ      TYPE      STATE  ADDR
lo0/v4       static    ok     192.68.0.1/8
net0/v4       static    ok     172.16.15.232/24
net1/v4       static    ok     172.16.16.220/24
net0/v6       addrconf  ok     fe80::203:baff:fe11:b115/10
lo0/v6       static    ok     ::1/128
net0/v6a      static    ok     2001:db8:3c4d:15:203:baff:fe11:b115/64
net1/v6       addrconf  ok     fe80::203:baff:fe11:b116/10
net1/v6a      static    ok     2001:db8:3c4d:16:203:baff:fe11:b116/64
```

Dans cet exemple, chaque interface configurée pour IPv6 possède maintenant deux adresses. L'entrée avec le nom d'objet d'adresse comme *interface/v6* indique l'adresse lien-local de l'interface. L'entrée avec le nom d'objet d'adresse comme *interface/v6add* indique une adresse globale IPv6. En plus de l'ID d'interface, cette adresse inclut le préfixe de site que vous avez configuré dans le fichier `/etc/ndpd.conf`. Notez que la désignation *v6add* est une chaîne définie de façon aléatoire. Vous pouvez définir d'autres chaînes pour constituer la seconde partie du nom d'objet d'adresse, à condition que l'*interface* reflète l'interface sur laquelle vous créez les adresses IPv6, par exemple `net0/mystring`, `net0/ipv6addr` et ainsi de suite.

- Voir aussi**
- Pour savoir comment configurer des tunnels à partir des routeurs identifiés dans la topologie de réseau IPv6, reportez-vous à “ [Administration des tunnels IP](#) ” du manuel “ [Administration des réseaux TCP/IP, d'IPMP et des tunnels IP dans Oracle Solaris 11.2](#) ”.
 - Pour obtenir des informations sur la configuration de commutateurs et de hubs sur votre réseau, reportez-vous à la documentation du fabricant.
 - Pour savoir comment améliorer la prise en charge d'IPv6 sur les serveurs, reportez-vous à “ [Administration d'interfaces compatibles IPv6 sur des serveurs](#) ” du manuel “ [Configuration et administration des composants réseau dans Oracle Solaris 11.2](#) ”.

Utilisation de Virtual Router Redundancy Protocol

La mise en place de matériel de secours pour remplacer les composants essentiels est un moyen d'accroître la fiabilité du réseau. Fournit un outil d'administration Oracle Solaris qui configure et gère l'utilisation du protocole de redondance de routeur virtuel (VRRP) pour assurer la haute disponibilité. VRRP est un protocole internet standard spécifié dans [RFC 5798](http://www.rfc-editor.org/rfc/rfc5798.txt) (<http://www.rfc-editor.org/rfc/rfc5798.txt>).

Fournit les propriétaires Oracle Solaris 11.2 permettant la prise en charge de la couche 3 VRRP routeurs sur VRRP InfiniBand VRRP des interfaces et de IP (IPMP) et la prise en charge existante VRRP améliorer dans les zones.

Remarque - Tout au long de ce chapitre, toutes les occurrences du terme VRRP Couche 2 (VRRP L2) font spécifiquement référence au VRPP Internet standard et toutes les occurrences du terme VRRP Couche 3 (VRRP L3) font référence au VRRP Couche 3 propriétaire d'Oracle Solaris.

Ce chapitre fournit une présentation de couche 2 VRRP dans et est la propriété Oracle Solaris couche 3.

Ce chapitre aborde les sujets suivants :

- “A propos de VRRP” à la page 27
- “Fonctionnement du protocole VRRP” à la page 28
- “A propos de la fonction VRRP de couche 3” à la page 30
- 3 “Comparaison de VRRP de couches 2 et 3” à la page 31
- “Limitations de routeur VRRP de couche 2 et 3” à la page 32

A propos de VRRP

De VRRP les adresses IP fournit une haute disponibilité, tels que ceux qui sont utilisés pour les routeurs ni des équilibreurs de charge. Les services qui utilisent sont également appelés

VRRP routeurs VRRP, elles exécutent les actions même si la gamme des services, tels que autre que l'équilibrage de charge. Pour plus d'informations sur la manière dont VRRP est utilisé avec l'équilibreur de charge afin d'assurer une haute disponibilité, reportez-vous au [Chapitre 7, Configuration ILB pour une haute disponibilité](#).

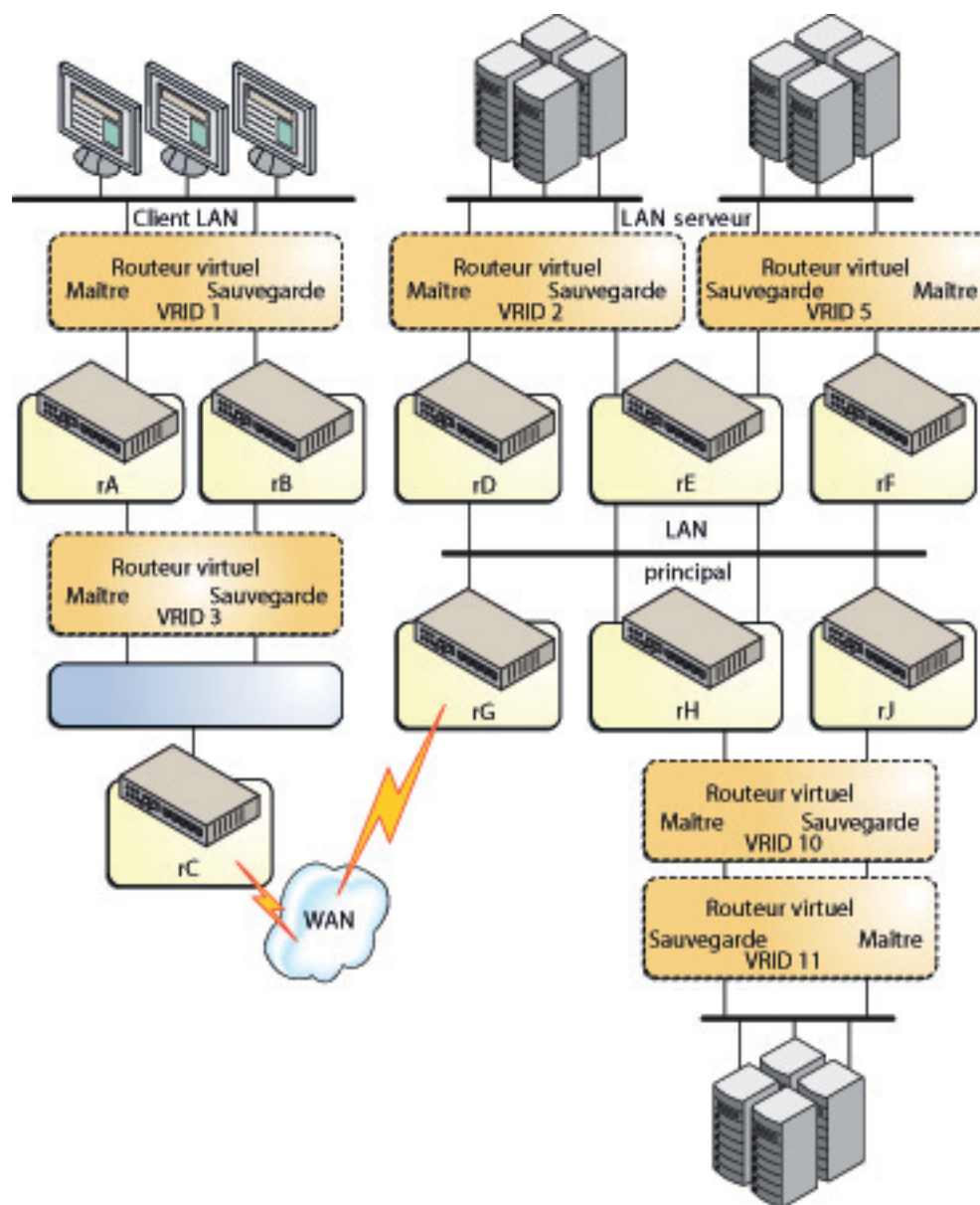
Un routeur VRRP est un routeur VRRP qui exécute le. Le service VRRP est exécuté sur chaque routeur VRRP et gère son état. Les routeurs un hôte peut disposer de plusieurs est configuré et sur lequel chaque routeur VRRP appartient à un routeur virtuel différent. Vous pouvez introduire des routeurs virtuels dans un réseau local (LAN) à l'aide de VRRP afin d'assurer la restauration après panne pour un routeur.

Fonctionnement du protocole VRRP

Veillez prendre note des éléments suivants pour le routeur VRRP :

- Nom de routeur : identifiant unique à l'échelle du système.
- ID de routeur virtuel (VRID) : numéro unique permettant d'identifier un routeur virtuel sur un segment donné au sein d'un réseau local.
- Adresse IP principale : adresse IP source de la publication VRRP.
- Adresse IP virtuelle (VRIP) : adresse IP associée à un VRID à partir de laquelle d'autres hôtes peuvent obtenir un service réseau et gérée par les instances VRRP appartenant à un VRID
- Routeur maître : instance VRRP qui assure la fonction de routage pour le routeur virtuel à un moment donné. Un seul *routeur maître* est actif à la fois pour un VRID donné. Le routeur maître contrôle la ou les adresse(s) IPv4 ou IPv6 associée(s) au routeur virtuel. Le routeur virtuel transfère les paquets envoyés à l'adresse IP du routeur maître.
- Routeur de secours : une instance VRRP pour un VRID actif mais pas en état maître est appelée un *routeur de secours*. N'importe quel nombre de routeurs de secours peut être défini pour un VRID. Un routeur de secours est prêt à prendre le rôle du routeur maître en cas de défaillance de ce dernier.
- Paramètres VRRP : priorité, intervalle de publication, mode de préemption et mode d'acceptation.
- Informations d'état et statistiques VRRP.

La configuration de partage de charge VRRP suivantes voit que le plusieurs VRID peuvent figurer sur une interface de routeur unique. Le texte d'accompagnement décrit les composants de VRRP utilisés dans la figure. Cette configuration de partage de charge VRRP montre que plusieurs VRID peuvent figurer sur une seule et même interface de routeur.

FIGURE 3-1 VRRP - Partage de chargement dans une configuration LAN

- Le routeur rA fait office de routeur maître pour le routeur virtuel VRID 1 et de routeur de secours pour VRID 3. Le routeur rA assure le routage des paquets adressés au VIP pour VRID 1, et il est prêt à prendre le rôle de routage pour VRID 3.
- Le routeur rB fait office de routeur maître pour le routeur virtuel VRID 3 et de routeur de secours pour VRID 1. Le routeur rB assure le routage des paquets adressés à l'adresse VIP de VRID 3 et il est prêt à assurer le routage pour VRID 1.
- Le routeur rC ne propose pas de fonctions VRRP mais il utilise l'adresse VIP de VRID 3 pour accéder au sous-réseau local du client.
- Le routeur rD fait office de routeur maître pour VRID 2. Le routeur rF fait office de routeur maître pour VRID 5. Le routeur rE est le routeur de secours pour ces deux VRID. En cas de défaillance de rD ou rF, rE devient le routeur maître pour ce VRID. Les routeurs rD et rF peuvent tomber en panne en même temps. Le routeur VRRP peut être un routeur maître pour un ou plusieurs deux VRID.
- Le routeur rG joue le rôle de passerelle WAN vers le réseau local principal. Tous les routeurs connectés au réseau principal partagent des informations de routage avec les routeurs du réseau WAN en utilisant un protocole de routage dynamique comme OSPF (Open Shortest Path First). VRRP n'est pas impliqué dans ces opérations, mais le routeur rC indique que le chemin d'accès au sous-réseau local du client passe par l'adresse VIP de VRID 3.
- Le routeur rH fait office de routeur maître pour VRID 10 et de routeur de secours pour VRID 11. De même, le routeur rJ fait office de routeur maître pour VRID 11 et de routeur de secours pour VRID 10.

A propos de la fonction VRRP de couche 3

La fonction de protocole propriétaire Virtual Router Redundancy Protocol Couche 3 (VRRP L3) d'Oracle Solaris évite d'avoir à configurer des adresses MAC virtuelles VRRP uniques pour les routeurs VRRP et assure ainsi une meilleure prise en charge des interfaces VRRP via IPMP et InfiniBand, ainsi que dans les zones. Le protocole VRRP L3 n'est pas conforme à la norme VRRP spécification. Au lieu d'utiliser une adresse MAC virtuelle VRRP unique parmi les routeurs dans le même routeur virtuel VRRP, l'implémentation utilise les L3 le protocole ARP (Address Resolution Protocol gratuitous les messages et ARP) NDP (Neighbor Discovery Protocol) des messages à régénérer la mise en correspondance entre les adresses MAC virtuel adresse IP du et le routeur VRRP maître en cours.

Fournit la couche 3 VRRP l'avantage de la prise en charge pour les interfaces InfiniBand sur IP (IPMP) et dans les zones et, ainsi que le format d'heure offrent une meilleure prise en charge évite d'avoir à créer la VNIC VRRP.

Comparaison de VRRP de couches 2 et 3

Le tableau suivant fournit une comparaison des VRRP et de couche 3 au niveau de la couche 2.

TABLEAU 3-1 Comparaison des VRRP couche 2 et la couche 3

Caractéristique	VRRP de couche 2	VRRP de couche 3
Création d'une VNIC VRRP	Vous devez créer une VNIC VRRP.	Vous n'avez pas à créer de VRRP VNIC VRRP adresse MAC car le virtuel qui est fourni par la carte réseau virtuelle VRRP n'est pas nécessaire.
La prise en charge d'IPMP	Non pris en charge.	Interfaces. Lorsqu'un routeur VRRP de couche 3 IPMP est créé par le biais d'une interface IP virtuel groupe sur chaque adresse, le routeur maître est associée à une adresse MAC IPMP de l'environnement d'initialisation actif existant sous-jacente IPMP interface en fonction de la stratégie. Si le basculement survient dans le ou le groupe IPMP les mappings L2, L3 gratuits sont annoncés par ARP ou NDP à l'aide des messages.
Prise en charge de zones	Il existe des problèmes pour exécuter plusieurs routeurs VRRP appartenant au même routeur virtuel dans différentes zones. Sur un système disposant de deux ou plusieurs routeurs qui partagent le même VRRP adresse MAC virtuelle VRRP, la classe intégrée de commutateur virtuel interrompt le flux normal du au VRRP routeur VRRP des paquets de publication. Pour plus d'informations, reportez-vous à "Limitations de routeur VRRP de couche 2 et 3" à la page 32.	Interfaces.
La prise en charge InfiniBand	Non pris en charge.	Interfaces.
Adresse MAC de routeur virtuel unique	Routeur virtuel requiert une adresse MAC unique. Les adresses IP virtuelles toujours la référence à la même adresse MAC virtuelle.	Pas nécessaire. Adresse MAC utilise sur lequel le la création du routeur VRRP. L'adresse MAC VRRP est différente d'une tous les routeurs qui sont dans le même routeur virtuel. La même adresse est associée au MAC virtuelle VRRP adresses IP cette L3 protégé par un routeur VRRP.
La configuration des adresses IP virtuelles VRRP	Configuration nécessaire.	Configuration nécessaire.
Redirections d'ICMP (Internet Control Message Protocol)	L2 peut être utilisée pour le groupe à l'autre VRRP est en cours d'exécution des routeurs. Lorsqu'un routeur VRRP L2 doit faire appel aux redirections ICMP, il vérifie l'adresse de destination MAC (adresse MAC virtuelle VRRP) des paquets qui doivent être réacheminés. A l'aide	Vous devez désactiver les réacheminements ICMP. Lorsque plusieurs routeurs VRRP sont créés sur la même interface, ils partagent la même adresse MAC. Par conséquent, le VRRP L3 ne peut pas déterminer l'adresse MAC de destination.

Caractéristique	VRRP de couche 2	VRRP de couche 3
	de l'adresse MAC de destination, le routeur VRRP L2 détermine le routeur virtuel auquel le paquet a été initialement envoyé. Par conséquent, le routeur VRRP L2 peut sélectionner l'adresse source et envoyer le message de redirection ICMP à la source.	
Sélection du routeur maître	La sélection du routeur maître n'a pas d'incidence pour l'hôte. Lorsqu'un routeur maître est modifié, le basculement entre l'hôte et le routeur identifie le nouveau port auquel envoyer le trafic à l'aide de sa fonction d'apprentissage MAC.	La sélection du routeur maître modifie le mappage de couche 2 des adresses IP virtuelles et un nouveau mappage doit être envoyé par les messages ARP ou NDP.
Durée de basculement	Normale.	Peut être plus long en raison de la condition requise par les messages ARP ou NDP lors de la sélection des modifications de routeur maître.

Limitations de routeur VRRP de couche 2 et 3

Couche 3 à la fois au niveau de la couche 2 et limite qui VRRP possèdent la même couche 2 vous devez configurer et de couche 3 les adresses IP virtuelles VRRP statiquement. Vous ne pouvez pas configurer automatiquement les adresses IP virtuelles VRRP en utilisant les deux outils de configuration automatique d'adresses IP `in.ndpd` pour la configuration automatique IPv6 et `dhcagent` pour la configuration DHCP (Dynamic Host Configuration Protocol). En outre, VRRP et de couche 3 au niveau de la couche 2 sont associés à des restrictions.

VRRP la fonction de la couche 2 présente les restrictions suivantes :

- **Prise en charge de zone en mode IP exclusif**

Quand un routeur VRRP est créé dans une zone d'IP exclusive, le service VRRP `svc:/network/vrrp/default` est activé automatiquement. Le service VRRP gère le routeur VRRP de cette zone spécifique. Toutefois, la prise en charge de zone IP exclusif est limitée comme suit :

- Etant donné qu'il n'est pas possible de créer une carte d'interface réseau virtuelle (VNIC, Virtual Network Interface Card) dans une zone non globale, vous devez tout d'abord créer la VNIC VRRP dans la zone globale. Par conséquent, commencez par créer la VNIC VRRP dans la zone globale, puis assignez la VNIC à la zone non globale où réside le routeur VRRP. Vous pouvez maintenant créer le routeur VRRP dans la zone non globale à l'aide de la commande `vrrpadm`.
- Oracle Solaris sur un système unique, vous ne pouvez pas créer deux routeurs VRRP dans des zones différentes partageant le même routeur virtuel. Oracle Solaris n'autorise pas la création de deux cartes réseau virtuelles avec la même MAC (Media Access Control) adresse.

- **Interopérations avec d'autres fonctionnalités du réseau**

- Le service VRRP L2 ne peut pas fonctionner sur une interface IPMP (fonctionnalité de chemins d'accès multiples sur réseau IP). VRRP nécessite en effet des adresses MAC

propres à VRRP alors que la fonctionnalité de chemins d'accès multiples sur réseau IP fonctionne totalement au niveau de la couche IP. Pour plus d'informations sur l'IPMP, reportez-vous au [Chapitre 2, “ IPMP d'administration sur ” du manuel “ Administration des réseaux TCP/IP, d'IPMP et des tunnels IP dans Oracle Solaris 11.2 ”](#).

VRRP peut être utilisé sur les groupements de liaisons dans les modes tronçon ou d'un groupement DLMP. Pour plus d'informations sur des agrégations, reportez-vous au [Chapitre 2, “ Configuration de la haute disponibilité à l'aide de groupements de liaisons ” du manuel “ Gestion des liaisons de données réseau dans Oracle Solaris 11.2 ”](#)

- Le service VRRP L2 ne peut pas fonctionner sur une IPoIB (IP over Infiniband) interface.

- **Prise en charge d'Ethernet via InfiniBand**

VRRP ne prend pas en charge la L2 Ethernet over InfiniBand (EoIB) interface. Routeur VRRP car chaque L2 virtuel unique est associé à une adresse, la valeur de MAC routeurs VRRP participant au même routeur virtuel doivent utiliser la même adresse MAC virtuel qui n'est pas pris en charge simultanément, par l'interface EoIB VRRP L3 dépasse cette limite lorsqu'il utilise une autre adresse MAC parmi les routeurs VRRP qui existent sur le même routeur virtuel.

Un VRRP de la couche 3 présente les restrictions suivantes :

- Les messages à l'aide ARP ou NDP gratuits risquent de se traduire par une au cours de la plus à l'intérieur du régime de la durée de basculement de routeur maître.

VRRP L3 utilise des messages ARP ou NDP pour envoyer les nouveaux mappages L2 ou L3 lors de la sélection des modifications du routeur maître. Cette exigence supplémentaire de l'utilisation de messages ARP ou NDP peut entraîner un plus long délai de basculement. Dans certains cas, si tous les messages annoncé ARP ou NDP gratuits sont perdues, l'opération peut prendre plus de temps à un hôte pour recevoir entrée ARP ou NDP il. Par conséquent, l'envoi de paquets sur le nouveau serveur maître routeur risque d'être retardé.

- Impossible de déterminer l'adresse MAC destination à l'aide de la redirection ICMP, car la même adresse MAC de destination est partagée par plusieurs routeurs.

Vous pouvez utiliser ICMP lorsque vous utilisez VRRP parmi un groupe de routeurs dans une topologie de réseau qui n'est pas symétrique. L'adresse source IPv4 ou IPv6 d'une redirection ICMPv4 ou ICMPv6 doit être l'adresse utilisée par la fin de l'hôte lorsqu'il a effectué la méthode de routage du prochain saut.

Lorsqu'un routeur VRRP L3 doit faire appel à des redirections ICMP, le routeur VRRP L3 vérifie l'adresse MAC de destination (adresse MAC virtuelle) des paquets qui doivent être réacheminés. La même adresse MAC de destination étant partagée par plusieurs routeurs créés sur la même interface, le routeur VRRP L3 ne peut pas déterminer l'adresse MAC de destination. Par conséquent, il peut s'avérer utile de désactiver la redirection ICMP lorsque vous utilisez des routeurs VRRP L3. Vous pouvez désactiver `send_redirects` grâce à la fonction de redirection ICMP public à l'aide du protocole IPv4 et IPv6 de la manière suivante :

```
# ipadm set-prop -m ipv4 -p send_redirects=off
```

- Il n'est pas possible de configurer les adresses IP virtuelles VRRP automatiquement par `in.ndpd` ou DHCP.

◆ ◆ ◆ 4 C H A P I T R E 4

Configuration et administration du protocole de redondance de routeur virtuel

Ce chapitre décrit les tâches de configuration des couches 2 et 3 de VRRP.

Ce chapitre aborde les sujets suivants :

- [“Planification d'une configuration VRRP” à la page 35](#)
- [“Installation de VRRP” à la page 36](#)
- [“Configuration VRRP” à la page 36](#)
- [“Cas d'utilisation : Configuration d'un routeur VRRP de couche 2” à la page 46](#)

Planification d'une configuration VRRP

La planification d'une couche 2 ou 3 de configuration VRRP implique les étapes suivantes :

1. Détermination du besoin de configurer un routeur VRRP de couche 2 ou de couche 3.
2. (pour le routeur VRRP de couche 2 uniquement) Création d'un VNIC VRRP. Pour plus d'informations, reportez-vous à la section [“Création d'une VNIC VRRP pour le VRRP de couche 2” à la page 37](#).

Vous pouvez créer automatiquement une VNIC VRRP en utilisant l'option -f de la commande `vrrpadm` quand vous créez le routeur VRRP L2.

3. Création d'un routeur VRRP Pour plus d'informations, reportez-vous à la section [“Création d'un routeur VRRP” à la page 38](#).
4. Configuration de l'adresse IP virtuelle pour le routeur VRRP. Pour plus d'informations, reportez-vous à la section [“Configuration de l'adresse IP virtuelle pour les routeurs VRRP de couches 2 et 3” à la page 40](#).

Vous pouvez configurer les adresses IP virtuelles à l'aide de l'option -a de la commande `vrrpadm`. Pour plus d'informations, reportez-vous à la section [“Création d'un routeur VRRP” à la page 38](#).

Installation de VRRP

Vous devez installer VRRP pour utiliser VRRP sur votre système.

▼ Installation de VRRP

1. **Connexion en tant qu'administrateur.**

Pour plus d'informations, reportez-vous à la section “ [A l'aide de vos droits administratifs attribués](#) ” du manuel “ [Sécurisation des utilisateurs et des processus dans Oracle Solaris 11.2](#) ”.

2. **Vérifiez si le package VRRP a été installé.**

```
# pkg info vrrp
```

3. **Installez le package VRRP s'il n'est pas installé.**

```
# pkg install vrrp
```

Configuration VRRP

Vous pouvez utiliser la commande `vrrpadm` pour configurer un routeur VRRP. Les résultats de toutes les sous-commandes de la commande `vrrpadm` sont persistants, à une exception près pour la commande `vrrpadm show-router`. Par exemple, le routeur VRRP créé à l'aide de la commande `vrrpadm create-router` sera conservé après les réinitialisations du système. Pour plus d'informations, reportez-vous à la page de manuel [vrrpadm\(1M\)](#).

Vous avez besoin d'une autorisation `solaris.network.vrrp` qui est incluse dans le profil Network Management, pour configurer le routeur VRRP.

Remarque - L'opération en lecture seule initiée par la commande `vrrpadm show-router` ne requiert pas d'autorisation `solaris.network.vrrp`.



Attention - Lorsque vous utilisez les VRRP fournis en standard avec le filtre IP regroupé d'Oracle Solaris, vous devez vérifier si le trafic IP entrant ou sortant pour le VRRP standard est autorisé pour l'adresse de multidiffusion 224.0.0.18/32, à l'aide de la commande `ipfstat -io`. Si le trafic n'est pas autorisé, les routeurs VRRP et maître seront en état MASTER. Par conséquent, vous devez ajouter les règles correspondantes à la configuration de filtre IP pour chaque routeur VRRP. Pour plus d'informations, reportez-vous à la section “[Dépannage des problèmes relatifs au VRRP et à Oracle Solaris Bundled IP Filter](#)” du manuel “[Dépannage des problèmes d'administration réseau dans Oracle Solaris 11.2](#)”.

Création d'une VNIC VRRP pour le VRRP de couche 2

Les VNIC sont des interfaces réseau virtuelles configurées sur un adaptateur réseau physique du système et sont des composants essentiels de la virtualisation de réseau. Une interface physique peut compter plusieurs cartes réseau virtuelles (VNIC). Pour plus d'informations sur VNIC, reportez-vous à la section “[Gestion de la virtualisation réseau et des ressources réseau dans Oracle Solaris 11.2](#)”.

Chaque routeur VRRP de couche 2 nécessite une VNIC VRRP spéciale. Utilisez la syntaxe de commande suivante.

```
# dladm create-vnic [-t] [-R root-dir] -l link [-m vrrp -V VRID -A \
{inet | inet6}] [-v VLAN-ID] [-p prop=value[,...]] VNIC
```

Cette commande crée une VNIC avec une adresse MAC de routeur virtuel définie par la spécification VRRP. Utilisez le type d'adresse VNIC `vrrp` pour spécifier le VRID et la famille d'adresse. La famille d'adresses est `inet` ou `inet6` qui fait référence à l'adresse IPv4 ou IPv6. Par exemple :

```
# dladm create-vnic -m vrrp -V 21 -A inet6 -l net0 vnic0
```

Pour plus d'informations, reportez-vous à la page de manuel [dladm\(1M\)](#)

Remarque - Vous pouvez également créer une VNIC VRRP à l'aide de l'option `-f` avec la commande `vrrpadm`. Pour plus d'informations, reportez-vous à la section “[Création d'un routeur VRRP](#)” à la page 38.

Création d'un routeur VRRP

La commande `vrrpadm create-router` crée un routeur VRRP de couche 2 ou de couche 3 avec le VRID et la famille d'adresse indiqués, ainsi que d'autres paramètres spécifiés. Pour plus d'informations, reportez-vous à la page de manuel [vrrpadm\(1M\)](#).

Pour créer un routeur VRRP, utilisez la syntaxe suivante :

```
# vrrpadm create-router [-T {l2 | l3}] [-f] -V VRID -I ifname \
-A [inet | inet6] [-a assoc-IPaddress] [-P primary-IPaddress] \
[-p priority] [-i adv-interval] [-o flags] router-name
```

<code>-T l2 l3</code>	Indique le type du routeur. Vous pouvez définir le type sur l'une des valeurs suivantes. L'inscription par défaut est <code>l2</code> . ■ <code>l2</code> : routeur VRRP de type couche 2 ■ <code>l3</code> : routeur VRRP de type couche 3
<code>-f</code>	(routeur VRRP de couche 2 uniquement) Indique la création de la VNIC VRRP avec un routeur VRRP de couche 2. Lorsque vous spécifiez l'option <code>-f</code> , la commande <code>vrrpadm</code> vérifie si la VNIC VRID avec le VRID et la famille d'adresses spécifiés existe. Une VNIC VRRP est créée uniquement si elle n'existe pas déjà. Le système génère le nom de la VNIC VRRP avec la convention de dénomination : <code>vrrp-VRID_ifname_v4 6</code> . L'option <code>-f</code> n'a aucun effet lorsque vous créez un routeur VRRP de couche 3.
<code>-V VRID</code>	Routeur virtuel VLAN qui définit l'identificateur que lorsqu'ils sont associés à la famille d'adresses.
<code>-I ifname</code>	L'interface sur laquelle le routeur VRRP est configuré. Pour un VRRP de couche 2, l'interface peut être une liaison physique, un VLAN ou un regroupement. Pour un VRRP de couche 3, l'interface peut également inclure une interface IPMP, une interface gérée par DHCP et une interface InfiniBand. Ce lien détermine LAN dans lequel ce VRRP est en cours d'exécution.
<code>-A [inet inet6]</code>	La famille d'adresses <code>inet</code> ou <code>inet6</code> , qui fait référence à l'adresse IPv4 ou IPv6.
<code>-a assoc-IPaddress</code>	Spécifie la liste séparée par des virgules des adresses IP. Vous pouvez spécifier l'adresse IP dans l'un des formats suivants : ■ <code>IP-address[/prefix-length]</code> ■ <code>hostname[/prefix-length]</code> ■ <code>linklocal</code>

Si vous indiquez `linklocal`, une adresse IPv6 lien-local vrrp est configurée en fonction du VRID du routeur virtuel associé. La forme `linklocal` s'applique uniquement aux routeurs IPv6 VRRP. Vous pouvez combiner l'option `-a` avec l'option `-f` de façon à ce que VNIC est créée et fixée automatiquement.

<code>-P primary-IPaddress</code>	Spécifie l'adresse IP VRRP principale utilisée pour envoyer la VRRP.
<code>-p priority</code>	Indiqué la priorité de la sélection utilisée pour les opérations de routeur VRRP. La valeur par défaut est 255. Le routeur avec la valeur de priorité la plus élevée est sélectionné en tant que routeur maître.
<code>-i adv-interval</code>	L'annonce, en millisecondes. La valeur par défaut est 1000.
<code>-o flags</code>	Les modes de préemption et d'acceptation du routeur VRRP. Les valeurs sont <code>preempt</code> ou <code>no_preempt</code> ou <code>accept</code> ou <code>no_accept</code> . Par défaut, les modes <code>pre-empt</code> et <code>accept</code> sont définis sur <code>preempt</code> et <code>accept</code> respectivement.
<code>router-name</code>	La valeur <code>router-name</code> est l'identifiant unique de ce routeur VRRP. Seuls les caractères alphanumériques (a-z, A-Z, 0-9) et le trait de soulignement ('_') sont autorisés dans un nom de routeur, dont la longueur est par ailleurs limitée à 31 caractères.

EXEMPLE 4-1 Création d'un routeur VRRP de couche 2

L'exemple suivant illustre la création d'un routeur sur une liaison de données `net0`.

```
# dladm create-vnic -m vrrp -V 12 -A inet -l net0 vnic1
# vrrpadm create-router -V 12 -A inet -p 100 -I net0 l2router1
# vrrpadm show-router l2router1
NAME      VRID  TYPE  IFNAME AF   PRIO ADV_INTV MODE  STATE  VNIC
l2router1 12    L2    net0  IPv4 100  1000  e-pa- BACK  vnic1
```

Un routeur VRRP L2 `l2router1` est créé sur la liaison de données `net0` avec une famille d'adresses IPv4 et VRID 12. Pour plus d'informations sur la commande `vrrpadm show-router`, reportez-vous à [“Affichage des configurations de routeur VRRP de couches 2 et 3” à la page 42](#).

EXEMPLE 4-2 Création d'un routeur VRRP de couche 3

L'exemple suivant montre comment créer un routeur VRRP L3 par le biais d'une interface IPMP nommée `ipmp0`.

```
# vrrpadm create-router -V 6 -I ipmp0 -A inet -T l3 l3router1
# vrrpadm show-router
NAME      VRID TYPE IFNAME AF   PRIO ADV_INTV MODE  STATE VNIC
l3router1 6    L3  ipmp0 IPv4 255 1000  eopa- INIT  --
```

Un routeur VRRP L3 `l3router1` est créé sur la l'interface IPMP `ipmp0` avec une famille d'adresses IPv4 et un VRID 6. Pour plus d'informations sur la commande `vrrpadm show-router`, reportez-vous à [“Affichage des configurations de routeur VRRP de couches 2 et 3” à la page 42](#).

Configuration de l'adresse IP virtuelle pour les routeurs VRRP de couches 2 et 3

Pour configurer l'adresse IP pour un routeur VRRP L2, vous devez configurer l'adresse IP virtuelle de type `vrrp` sur la VNIC VRRP qui lui est associée.

Pour configurer l'adresse IP virtuelle d'un routeur VRRP L3, vous devez utiliser une adresse IP de type `vrrp` sur la même interface IP où est configuré le routeur VRRP L3.

Remarque - Pour configurer une adresse IPv6, vous devez avoir créé des VNIC VRRP ou le routeur L3 VRRP en indiquant la famille d'adresses du routeur sous la forme `inet6`.

Pour configurer une adresse IP virtuelle pour un routeur VRRP, utilisez la syntaxe suivante :

```
# ipadm create-addr [-t] -T vrrp [-a local=addr[/prefix-length]] \
  [-n router-name]... addr-obj | interface
```

`-t` Indique que l'adresse configurée est temporaire et que les modifications s'appliquent uniquement à la configuration active.

`-T vrrp` Indique que l'adresse configurée est de type `vrrp`.

`-n router-name` L'option `-n router-name` est facultative pour un routeur VRRP L2, car le nom du routeur VRRP peut être obtenu à partir de l'interface VNIC VRRP sur laquelle les adresses IP sont configurées.

Pour plus d'informations, reportez-vous à la page de manuel [ipadm\(1M\)](#).

Remarque - Vous pouvez également configurer les adresses IP virtuelles à l'aide de l'option `-a` avec la commande `vrrpadm`. Pour plus d'informations, reportez-vous à la section [“Création d'un routeur VRRP” à la page 38](#).

EXEMPLE 4-3 La configuration L2 l'adresse IP virtuelle VRRP pour un routeur

Vous pouvez utiliser l'adresse IP `vrrp` pour configurer les adresses IP virtuelles d'un routeur VRRP L2. L'exemple ci-dessous montre comment créer l'adresse IP virtuelle pour `l2router1`.

```
# ipadm create-ip vrrp_vnic1
# ipadm create-addr -T vrrp -n l2router1 -a 192.168.82.8/24 vrrp_vnic1/vaddr1
```

L'exemple suivant montre comment créer une adresse IP IPv6 lien-local `vrrp` pour `V6vrrp_vnic1/vaddr1`.

```
# ipadm create-ip V6vrrp_vnic1
# ipadm create-addr -T vrrp V6vrrp_vnic1/vaddr1
```

Pour configurer l'adresse de type IP IPv6 link-local `vrrp` pour un routeur VRRP, il est inutile d'entrer l'adresse locale. Une adresse de type IP IPv6 link-local `vrrp` est créée en fonction du VRID du routeur VRRP associé.

EXEMPLE 4-4 La configuration du L3 l'adresse IP virtuelle VRRP pour un routeur

L'exemple ci-dessous montre comment configurer l'adresse IP virtuelle pour `l3router1`.

```
# ipadm create-ip ipmp0
# ipadm create-addr -T vrrp -n l3router1 -a 172.16.82.8/24 ipmp0/vaddr1
```

L'exemple suivant montre comment configurer une adresse IP type IPv6 link-local `vrrp` pour le routeur VRRP L3 `l3V6router1`.

```
# ipadm create-ip ipmp1
# ipadm create-addr -T vrrp -n l3V6router1 ipmp1/vaddr0
```

Activation et désactivation d'un routeur VRRP

Un routeur VRRP est activé par défaut lorsque vous le créez. Vous pouvez désactiver un routeur VRRP et la réactiver. La liaison de données sous-jacente sur laquelle est créé le routeur VRRP (spécifiée par l'option `-l` lors de la création du routeur par le biais de la commande `vrrpadm create-router`) et sa carte réseau virtuelle VRRP doivent exister lors de l'activation du routeur. Sans quoi, l'opération d'activation échoue. Pour un routeur VRRP de couche 2, si la VNIC VRRP du routeur n'existe pas, le routeur n'est pas activé. Respectez la syntaxe suivante :

```
# vrrpadm enable-router router-name
```

Il est parfois nécessaire de désactiver temporairement un routeur VRRP pour apporter des modifications à la configuration, puis de réactiver le routeur. Pour la désactivation d'un routeur, la syntaxe est la suivante :

```
# vrrpadm disable-router router-name
```

Modification d'un routeur VRRP

La commande `vrrpadm modify-router` modifie la configuration d'un routeur VRRP spécifié. Vous pouvez modifier la priorité, l'intervalle de publication, le mode pre-empt et le mode accept du routeur. Respectez la syntaxe suivante :

```
# vrrpadm modify-router [-p priority] [-i adv-interval] [-o flags] router-name
```

Affichage des configurations de routeur VRRP de couches 2 et 3

La commande `vrrpadm show-router` affiche la configuration et l'état d'un routeur VRRP spécifié. Pour plus d'informations, reportez-vous à la page de manuel [vrrpadm\(1M\)](#). Respectez la syntaxe suivante :

```
# vrrpadm show-router [-P | -x] [-p] [-o field[,...]] [router-name]
```

EXEMPLE 4-5 Affichage de la configuration du routeur VRRP de couche 2

Les exemples suivants illustrent la sortie de commande `vrrpadm show-router`.

```
# vrrpadm show-router vrrp1
NAME VRID TYPE IFNAME AF PRIO ADV_INTV MODE STATE VNIC
vrrp1 1 L2 net1 IPv4 100 1000 e-pa- BACK vnic1
```

NAME Le nom du routeur VRRP.

VRID VRID du VRRP du routeur.

TYPE Le type de routeur VRRP ou qui est soit L2, L3.

IFNAME L'interface sur laquelle le routeur VRRP est configuré. Pour un routeur VRRP de couche 2, l'interface peut être une interface physique Ethernet, un VLAN ou un regroupement.

AF La famille d'adresses du routeur VRRP. Il peut s'agir d'IPv4 ou IPv6.

PRIO Le niveau de priorité du routeur VRRP, qui est utilisé pour les opérations de sélection.

ADV_INTV	L'intervalle de publication affiché en millisecondes.
MODE	Un ensemble d'indicateurs associés au routeur VRRP et qui comprennent les valeurs possibles suivantes : ■ e : indique que le routeur est activé. ■ p : indique que le mode est preempt. ■ a : indique que le mode est accept. ■ o : indique que le routeur est le propriétaire de l'adresse virtuelle.
STATE	L'état en cours du routeur VRRP. Les valeurs possibles sont INIT (initialiser), BACK (sauvegarde) et MAST (maître).

Dans cet exemple, des informations sur le routeur VRRP `vrrp1` spécifié s'affichent.

```
# vrrpadm show-router -x vrrp1
NAME STATE PRV_STAT STAT_LAST VNIC PRIMARY_IP VIRTUAL_IPS
vrrp1 BACK MAST 1m17s vnic1 10.0.0.100 10.0.0.1
```

PRV_STAT	L'état précédent du routeur VRRP.
STAT_LAST	Temps écoulé depuis la dernière transition d'état.
PRIMARY_IP	L'adresse IP principale sélectionnée par le routeur VRRP.
VIRTUAL_IPS	Les adresses IP virtuelles configurées sur le routeur VRRP.

Dans cet exemple, des informations supplémentaires sur le routeur, telles que l'adresse IP principale sélectionnée par le routeur VRRP, l'adresse IP virtuelle configurée sur le routeur VRRP et l'état précédent du routeur VRRP sont affichées.

```
# vrrpadm show-router -P vrrp1
NAME PEER P_PRIO P_INTV P_ADV_LAST M_DOWN_INTV
vrrp1 10.0.0.123 120 1000 0.313s 3609
```

peer	L'adresse IP principale du routeur VRRP homologue.
P_PRIO	Le niveau de priorité du routeur VRRP homologue, qui fait partie de l'annonce en provenance de l'homologue.
P_INTV	L'intervalle de publication (en millisecondes), qui fait partie des publications en provenance de l'homologue.
P_ADV_LAST	Temps depuis la dernière publication reçue de l'homologue.
M_DOWN_INTV	Intervalle de temps (en millisecondes) au bout duquel le routeur maître peut être déclaré comme étant arrêté.

Vous devez utiliser l'option `-P` uniquement lorsque le routeur VRRP est en état de sauvegarde.

EXEMPLE 4-6 Affichage du routeur VRRP de couche 3 sur un système

```
# vrrpadm show-router
NAME  VRID  TYPE  IFNAME  AF    PRIO  ADV_INTV  MODE  STATE  VNIC
l3vr1 12    L3    net1    IPv6  255    1000      eopa-  INIT   -
```

Dans cet exemple, le routeur VRRP l3vr1 L3 est configuré sur l'interface net1.

Affichage des adresses IP associées à des routeurs VRRP

Vous pouvez afficher l'adresse IP associée avec un routeur VRRP à l'aide de la commande `ipadm show-addr`. Le champ `ROUTER` de la sortie de la commande `ipadm show-addr` affiche le nom du routeur VRRP spécifique qui est associé à une adresse IP type `vrrp`.

Pour l'adresse IP type `vrrp` d'un VRRP L2, le nom du le routeur VRRP est dérivé de la VNIC VRRP pour laquelle l'adresse IP est configurée. Si vous exécutez la commande `ipadm show-addr` avant de créer le routeur L2 pour une VNIC VRRP, le champ `ROUTER` affiche `?`. Pour l'adresse IP type `vrrp` d'un VRRP L3, le champ `ROUTER` affiche toujours le nom de routeur spécifié. Pour les autres types d'adresses IP, le champ `ROUTER` n'est pas applicable, et `--` s'affiche.

EXEMPLE 4-7 Affichage des adresses IP associées à des routeurs VRRP

```
# ipadm show-addr -o addrobj,type,vrrp-router,addr
ADDROBJ          TYPE  VRRP-ROUTER  ADDR
lo0/v4           static --          127.0.0.1/8
net1/p1           static --          192.168.11.10/24
net1/v1           vrrp   l3router1    192.168.81.8/24
vrrp_vnic1/vaddr1 vrrp   l2router1    192.168.82.8/24
lo0/v6           static --          ::1/128
```

Dans cet exemple, l3router1 est associé avec l'adresse IP type `vrrp` 192.168.81.8/24, et l2router1 est associé avec l'adresse IP type `vrrp` 192.168.82.8/24.

La sortie affiche les informations suivantes :

ADDROBJ Le nom de l'objet d'adresse.

TYPE Le type de l'objet d'adresse, qui peut être l'un des suivants :

- `from-gz`
- `static`

- dhcp
- addrconf
- vrrp

VRRP-ROUTER	Le nom du routeur VRRP.
ADDR	Adresse IPv4 ou IPv6 numérique.

Suppression d'un routeur VRRP

La commande `vrrpadm delete-router` supprime le routeur VRRP spécifié. Respectez la syntaxe suivante :

```
# vrrpadm delete-router router-name
```

Remarque - La VNIC VRRP, l'adresse IP type `vrrp` et l'adresse IP principale créées à l'aide des options `-f`, `-a`, `-P` de la commande `vrrpadm create-router` respectivement ne sont pas supprimées suite à la commande `vrrpadm delete-router`. Vous devez les supprimer de manière explicitement par le biais des commandes correspondantes `ipadm` et `dladm`.

Contrôle des messages Gratuitous ARP et NDP

Remplace le signe plus (+) lorsqu'une copie de sauvegarde VRRP routeur VRRP maître définit un routeur, l'indicateur IP virtuel sur toutes les adresses associées à l'IP de routeur maître. par conséquent, les adresses IP sont protégés. S'il n'y a pas de conflit concernant les adresses IP virtuelles, plusieurs messages de type ARP gratuit et "neighbour advertisement" sont envoyés pour publier le nouveau mappage entre l'adresse IP virtuelle et l'adresse MAC du nouveau maître.

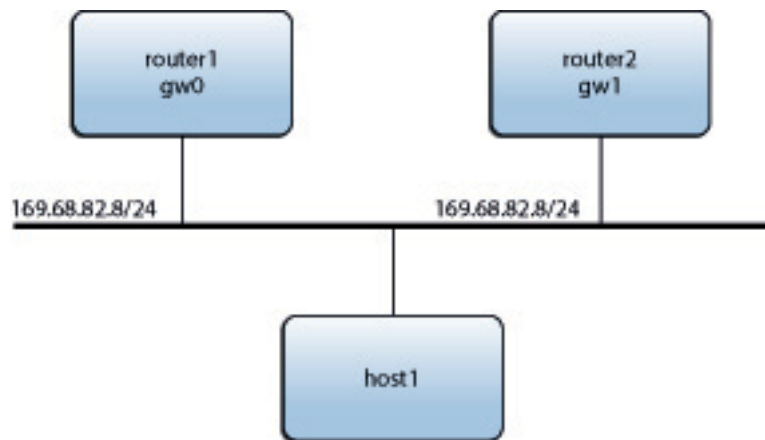
Pour contrôler le nombre de messages envoyés et l'intervalle entre l'annonce de messages, vous pouvez les propriétés de protocole IP dans le tableau ci-dessous :

- `arp_publish_count`
- `arp_publish_interval`
- `ndp_unsolicit_count`
- `ndp_unsolicit_interval`

Pour plus d'informations sur les propriétés de protocole IP, reportez-vous à [“ Paramètres réglables IP liés à la détection des adresses dupliquées ”](#) du manuel [“ Manuel de référence des paramètres réglables d'Oracle Solaris 11.2 ”](#).

Cas d'utilisation : Configuration d'un routeur VRRP de couche 2

Le schéma suivant illustre une configuration VRRP standard.



Dans cet exemple, l'adresse IP 10.68.82.8 est configurée en tant que passerelle par défaut pour l'hôte host1. Il s'agit de l'adresse IP virtuelle protégée par le routeur virtuel composé de deux routeurs VRRP, à savoir router1 et router2. A tout moment, un seul des deux routeurs fait office de routeur maître, en supposant que les responsabilités du routeur virtuel et le processus de transfert de paquets en provenance de host1.

Partons du principe que le VRID du routeur virtuel est 12. Les exemples ci-dessous illustrent les commandes permettant d'établir la configuration VRRP précédente sur les routeurs router1 et router2, router1 étant propriétaire de l'adresse IP virtuelle 10.68.82.8 dont la priorité correspond à la valeur par défaut (255), et router2 étant le routeur de secours dont la priorité est fixée à 100.

Pour plus d'informations sur les commandes utilisées pour la configuration de VRRP, reportez-vous aux pages de manuel [vrrpadm\(1M\)](#), [dladm\(1M\)](#) et [ipadm\(1M\)](#).

Pour router1 :

1. Installez le package VRRP.

```
# pkg install vrrp
```

2. Créez l'a VNIC vnic0 sur net0 avec la valeur VRID 12.

```
# dladm create-vnic -m vrrp -V 12 -A inet -l net0 vnic0
```

3. Créez un routeur VRRP vrrp1 via net0.

```
# vrrpadm create-router -V 12 -A inet -I net0 vrrp1
```

4. Configurez les interfaces IP vnic0 et net0.

```
# ipadm create-ip vnic0
```

```
# ipadm create-addr -T vrrp -a 10.68.82.8/24 vnic0/router1
```

```
# ipadm create-ip net0
```

```
# ipadm create-addr -T static -a 10.68.82.100/24 net0/router1
```

5. Affichez les informations de routeur pour vrrp1.

```
# vrrpadm show-router -x vrrp1
```

NAME	STATE	PRV_STAT	STAT_LAST	VNIC	PRIMARY_IP	VIRTUAL_IPS
vrrp1	MASTER	INIT	14.444s	vnic0	10.68.82.100	10.68.82.8

De même, pour router2 :

1. Créez la VNIC vnic1 sur net1 avec la valeur VRID définie sur 12.

```
# dladm create-vnic -m vrrp -V 12 -A inet -l net1 vnic1
```

2. Créez un routeur VRRP vrrp2 sur net1.

```
# vrrpadm create-router -V 12 -A inet -I net1 -p 100 vrrp2
```

3. Configurez les interfaces IP vnic1 et net1.

```
# ipadm create-ip vnic1
```

```
# ipadm create-addr -T vrrp -a 10.68.82.8/24 vnic1/router2
```

```
# ipadm create-ip net1
```

```
# ipadm create-addr -T static -a 10.68.82.101/24 net1/router2
```

4. Affichez les informations de routeur pour vrrp2.

```
# vrrpadm show-router -x vrrp2
```

NAME	STATE	PRV_STAT	STAT_LAST	VNIC	PRIMARY_IP	VIRTUAL_IPS
vrrp2	BACKUP	INIT	2m32s	vnic1	10.68.82.101	10.68.82.8

En utilisant la configuration de router1 comme exemple, vous devez configurer au moins une adresse IP sur gw0. L'adresse IP du router1 correspond à l'adresse IP principale utilisée pour envoyer les paquets de publication VRRP.

```
# vrrpadm show-router -x vrrp1
```

NAME	STATE	PRV_STAT	STAT_LAST	VNIC	PRIMARY_IP	VIRTUAL_IPS
vrrp1	MASTER	INIT	14.444s	vnic1	10.68.82.100	10.68.82.8

Présentation d'un équilibreur de charge intégré

Ce chapitre décrit les composants d'ILB et les modes de fonctionnement d'ILB, telles que DSR (Direct Server Return) et le mode Network Address Translator (NAT).

Ce chapitre aborde les sujets suivants :

- [“Composants ILB” à la page 49](#)
- [“Modes de fonctionnement d'ILB” à la page 50](#)
- [“Fonctionnement de l'équilibreur de charge intégré” à la page 54](#)

Pour plus d'informations sur l'ILB, reportez-vous à la section [“Présentation de l'équilibreur de charge intégré” à la page 13](#).

Composants ILB

ILB est géré par l'utilitaire de gestion des services SMF (Service Management Facility) `svc:/network/loadbalancer/ilb:default`. Pour plus d'informations sur SMF, reportez-vous à la section [“Gestion des services système dans Oracle Solaris 11.2”](#). Les trois composants principaux d'ILB sont les suivants :

- CLI (interface de ligne de commande) `ilbadm` : vous pouvez configurer des règles d'équilibrage de la charge, effectuer des vérifications de l'état en option et consulter des statistiques dans cette interface de ligne de commande.
- Bibliothèque de configuration `libilb` : `ilbadm` et des applications tierces peuvent effectuer des opérations d'administration de l'équilibreur de charge intégré par le biais de la fonctionnalité implémentée dans `libilb`.
- Démon `ilbd` : il réalise les tâches suivantes :
 - Il gère la configuration persistante d'une réinitialisation ou d'une mise à jour de package à l'autre
 - Il offre un accès série au module de noyau ILB en traitant les données de configuration et en les envoyant au module de noyau ILB en vue de leur exécution.
 - Il effectue des vérifications de l'état des serveurs et envoie les résultats au module du noyau ILB afin d'optimiser la répartition de la charge.

Modes de fonctionnement d'ILB

ILB prend en charge les modes DSR et NAT sans état pour IPv4 et IPv6, par le biais de topologies à une ou deux branches.

Mode Direct Server Return (DSR)

En mode DSR, ILB répartit les demandes entrantes sur les serveurs backend. Le trafic de retour en provenance de ces serveurs vers les clients contourne ILB. Dans ce cas, la réponse du serveur backend envoyée au client est acheminée par le biais du système qui exécute ILB. L'implémentation actuelle de l'équilibreur de charge intégré en mode DSR n'offre pas de fonction de suivi des connexions TCP (sans conservation de statut). Avec le mode DSR sans état, ILB n'enregistre pas les informations d'état des paquets traités, excepté à des fins statistiques. Ce mode étant sans état, l'opération est comparable au transfert IP normal. Ce mode est particulièrement adapté aux protocoles sans connexion.

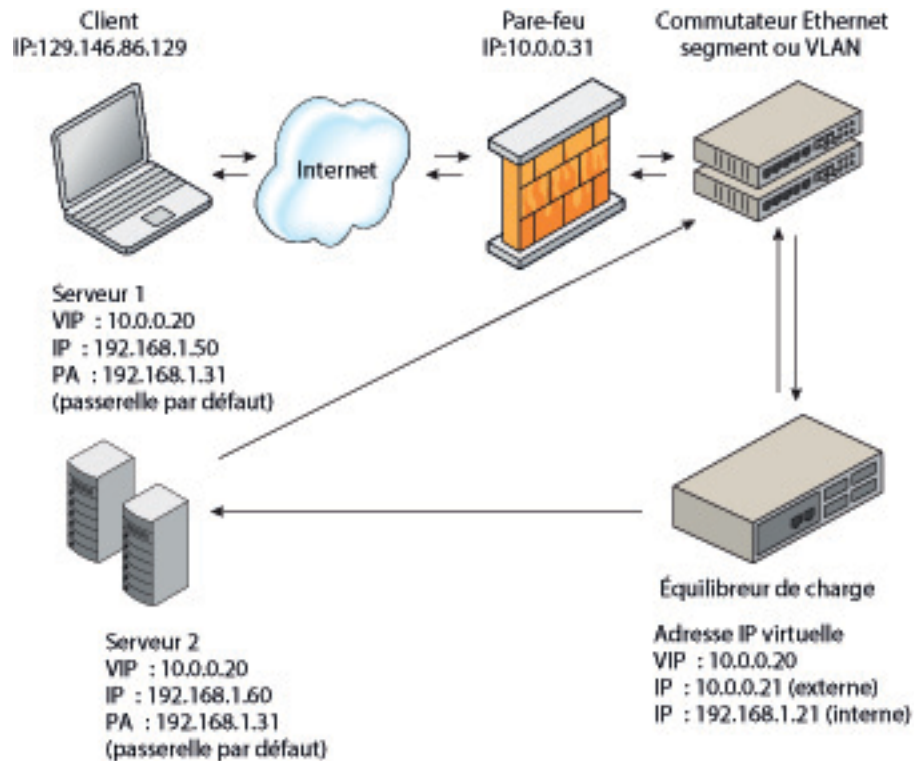
Avantages :

- Cette topologie offre de meilleures performances que le mode NAT, car seule l'adresse MAC de destination des paquets est modifiée, et les serveurs répondent directement aux clients.
- Vous bénéficiez d'une transparence totale entre le serveur et le client. Les serveurs perçoivent une connexion directement à partir de l'adresse IP client et répondent au client via la passerelle par défaut.

inconvénients :

- Le serveur backend doit répondre à la fois aux demandes envoyées à sa propre adresse IP (vérifications de l'état) et à l'adresse IP virtuelle (trafic avec équilibrage de la charge).
- Ce mode fonctionnant sans état, l'ajout ou la suppression de serveurs entraîne l'interruption des connexions.

La figure ci-dessous illustre l'implémentation d'ILB avec la topologie DSR.

FIGURE 5-1 Topologie DSR (Direct Server Return)

Dans ce schéma, les deux serveurs backend sont situés dans le même sous-réseau (192.168.1.0/24) que l'équilibreur de charge intégré. Ces serveurs sont également connectés au routeur de manière à pouvoir répondre directement aux clients après avoir reçu une demande transférée par ILB.

Mode NAT (Network Address Translator)

ILB utilise la topologie NAT en mode autonome exclusivement à des fins d'équilibrage de la charge. Dans ce mode, l'équilibreur de charge intégré réécrit les informations d'en-tête, et gère aussi bien le trafic entrant que le trafic sortant. ILB fonctionne en mode Half-NAT ou Full-NAT. Toutefois, la topologie Full-NAT réécrit également l'adresse IP source de sorte que le serveur considère que toutes les connexions proviennent de l'équilibreur de charge. Le mode NAT offre une fonction de suivi des connexions TCP (avec conservation de statut). Le mode NAT renforce

la sécurité et s'adapte mieux au trafic HTTP (Hypertext Transfer Protocol) et SSL (Secure Sockets Layer).

avantages :

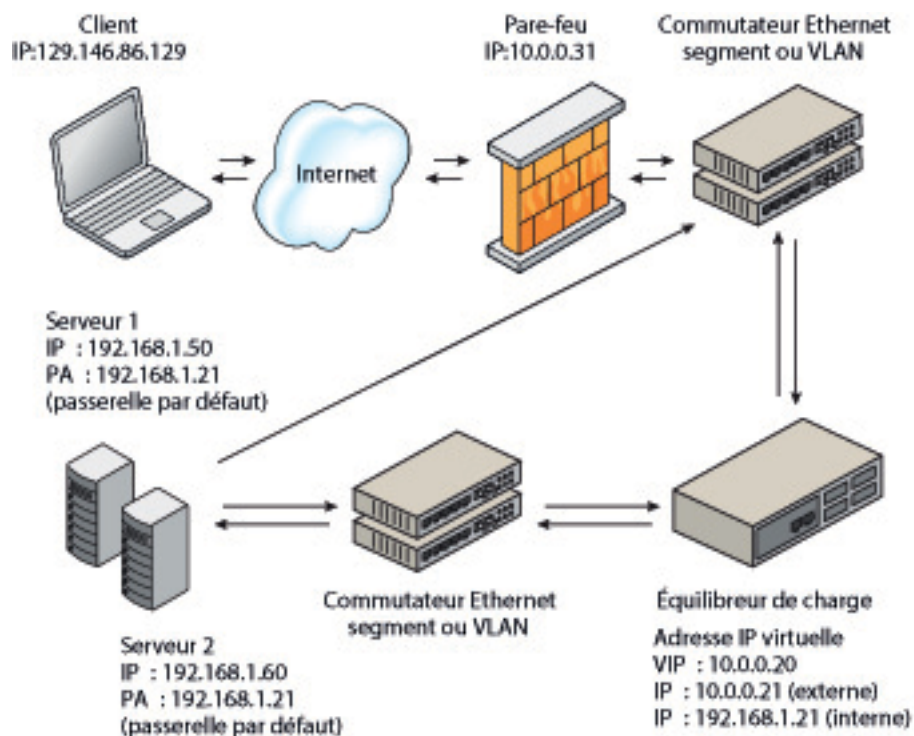
- Cette topologie fonctionne avec tous les serveurs backend en modifiant la passerelle par défaut pour pointer vers l'équilibreur de charge
- L'ajout ou la suppression de serveurs peut avoir lieu sans interruption des connexions du fait que l'équilibreur de charge maintient l'état de connexion

inconvénients :

- Ce mode est moins rapide que la topologie DSR car le traitement implique la manipulation de l'en-tête IP et les serveurs envoient des réponses à l'équilibreur de charge
- Tous les serveurs backend doivent utiliser l'équilibreur de charge en tant que passerelle par défaut

L'implémentation générale de la topologie NAT s'affiche comme illustré dans la figure suivante.

FIGURE 5-2 Topologie NAT (Network Address Translation)



Dans ce cas, toutes les demandes envoyées à l'adresse IP virtuelle passent par l'équilibreur de charge intégré puis sont transférées aux serveurs backend. En mode NAT, toutes les réponses provenant des serveurs backend passent par ILB.



Attention - Le chemin d'accès au code NAT implémenté dans l'équilibreur de charge intégré ne correspond pas à celui de la fonction de filtre IP d'Oracle Solaris. *N'utilisez pas les deux chemins d'accès au code simultanément.*

L'équilibrage Half-NAT Mode de chargement

En mode Half-NAT, ILB réécrit uniquement l'adresse IP de destination dans l'en-tête des paquets. Si vous mettez en oeuvre l'implémentation Half-NAT, vous ne pouvez pas vous connecter à l'adresse IP virtuelle (VIP) du service à partir du sous-réseau sur lequel réside le serveur. Le tableau suivant répertorie les adresses IP des paquets circulant entre un client et ILB, puis entre ILB et un serveur backend.

TABLEAU 5-1 Flux de demande et de réponse dans la topologie Half-NAT lorsque le serveur et le client résident sur différents réseaux

Flux de demande	Adresse IP source	Adresse IP de destination
1. Client -> ILB	Client	VIP d'ILB
2. ILB -> Serveur	Client	Serveur
Flux de réponse		
3. Serveur -> ILB	Serveur	Client
4. ILB -> Client	VIP d'ILB	Client

Si vous connectez le système client sur le même réseau que celui des serveurs, le serveur souhaité répond directement au client et la quatrième étape du processus dans la table n'ont pas lieu. Aussi, l'adresse IP source de la réponse que le serveur envoie au client n'est pas valide. Lorsque le client envoie une demande de connexion à l'équilibreur de charge, la réponse parvient au serveur concerné. A ce stade, la pile IP du client abandonne correctement toutes les réponses. Dans ce cas, les flux de demande et de réponse se présentent comme illustré dans le tableau suivant.

TABLEAU 5-2 Flux de demande et de réponse dans la topologie Half-NAT lorsque le serveur et le client résident sur le même réseau

Flux de demande	Adresse IP source	Adresse IP de destination
1. Client -> ILB	Client	VIP d'ILB
2. ILB -> Serveur	Client	Serveur
Flux de réponse		
3. Serveur -> Client	Serveur	Client

Mode d'équilibreur de charge Full-NAT

En mode Full-NAT, les adresses IP source et de destination sont réécrites pour garantir que le trafic passe par l'équilibreur de charge dans les deux directions. La topologie Full-NAT permet de se connecter à l'adresse VIP à partir du sous-réseau sur lequel sont situés les serveurs.

Le tableau suivant répertorie les adresses IP des paquets circulant entre un client et ILB, puis entre ILB et un serveur backend dans la topologie Full-NAT. Aucun routage par défaut par le biais de l'équilibreur de charge intégré n'est requis dans ces serveurs. Mais notez que l'administrateur doit impérativement réserver une adresse IP ou une plage d'adresses qui feront office d'adresses source de l'équilibreur de charge intégré dans le cadre des communications avec les serveurs backend dans une topologie Full-NAT. Partons du principe que les adresses utilisées appartiennent au sous-réseau C. Dans ce scénario, ILB se comporte comme un proxy.

TABLEAU 5-3 Flux de demande et de réponse dans la topologie Full-NAT

Flux de demande	Adresse IP source	Adresse IP de destination
1. Client -> ILB	Client	VIP d'ILB
2. ILB -> Serveur	Adresse de l'interface de l'équilibreur de charge (sous-réseau C)	Serveur
Flux de réponse		
3. Serveur -> ILB	Serveur	Adresse de l'interface de l'équilibreur de charge intégré (sous-réseau C)
4. ILB -> Client	VIP d'ILB	Client

Fonctionnement de l'équilibreur de charge intégré

Cette section décrit le processus d'ILB, qui implique le traitement d'une demande d'un client vers VIP, le transfert de la demande à un serveur backend et le traitement de la réponse.

Le traitement de paquets client-serveur implique les étapes suivantes :

1. ILB reçoit une demande entrante que le client envoie à une adresse IP virtuelle. Il l'analyse par rapport aux règles d'équilibrage de la charge configurées.
2. Si une règle d'équilibrage de la charge s'applique, ILB met en oeuvre un algorithme pour transférer la demande à un serveur backend en fonction de la topologie en place.
 - En mode DSR, ILB remplace l'en-tête MAC de la demande entrante par l'en-tête MAC du serveur d'arrière-plan sélectionné.
 - En mode Half-NAT, ILB remplace l'adresse IP cible et le numéro du port pour le protocole de transfert de la demande entrante par ceux du serveur d'arrière-plan sélectionné.

- En mode Full-NAT, ILB remplace l'adresse IP source et le numéro de port du protocole de transport de la demande entrante par l'adresse NAT source de la règle d'équilibrage de la charge. L'équilibreur de charge intégré remplace également l'adresse IP de destination et le numéro de port du protocole de transport de la demande entrante par ceux du serveur backend sélectionné.
- 3. L'équilibreur de charge intégré transfère la demande entrante modifiée au serveur backend sélectionné.

Le traitement de paquets serveur-client implique les étapes suivantes :

1. Le serveur backend envoie une réponse à ILB à la réception d'une demande entrante en provenance du client.
2. Le comportement de l'équilibreur de charge intégré à la réception d'une réponse du serveur backend dépend du mode de fonctionnement.
 - En mode DSR, la réponse du serveur backend contourne ILB et parvient directement au client. Cela étant, si l'équilibreur de charge intégré fait office de routeur, la réponse du serveur backend envoyée au client est acheminée par le biais du système qui exécute ILB.
 - En mode Half-NAT ou Full-NAT, l'équilibreur de charge intégré fait correspondre la réponse du serveur backend à la demande entrante, puis remplace l'adresse IP et le numéro de port du protocole de transport modifiés par ceux de la demande entrante d'origine. ILB transfère ensuite la réponse au client.

Configuration et gestion de l'équilibreur de charge intégré

ILB est configuré sur les couches L3 et L4 de la pile de protocole réseau. Ce chapitre décrit les tâches pour l'installation, l'activation ou la désactivation d'ILB, la définition de groupes de serveurs ILB serveurs back-end dans ILB, l'importation et l'exportation de configurations ILB à l'aide de la commande `ilbadm`. Pour plus d'informations, reportez-vous à la page de manuel [ilbadm\(1M\)](#). Pour plus d'informations sur la configuration de vos ILB pour la haute disponibilité, reportez-vous au [Chapitre 7, Configuration ILB pour une haute disponibilité](#).

Pour plus d'informations sur le déploiement de l'équilibreur de charge intégré Oracle Solaris, reportez-vous à la section [Deploying the Oracle Solaris Integrated Load Balancer in 60 Minutes](#) (<http://www.oracle.com/technetwork/systems/hands-on-labs/hol-deploy-ilb-60mmin-2137812.html>). Pour plus d'informations sur l'ajout à votre application de la haute disponibilité à l'aide d'Oracle Solaris Zones et de l'équilibreur de charge intégré à Oracle Solaris, reportez-vous à la section [How to Set Up a Load-Balanced Application Across Two Oracle Solaris Zones](#) (<http://www.oracle.com/technetwork/articles/servers-storage-admin/loadbalancedapp-1653020.html>).

Ce chapitre aborde les sujets suivants :

- “Installation d'ILB” à la page 57
- “Configuration d'ILB à l'aide de la CLI” à la page 58
- “Activation ou désactivation d'ILB” à la page 59
- “Gestion d'un ILB” à la page 60
- “Cas d'utilisation : configuration d'un ILB” à la page 70
- “Affichage des statistiques ILB” à la page 71
- “Importation et exportation des configurations” à la page 74

Installation d'ILB

ILB nécessite des installations de noyau et de l'utilisateur. L'installation du noyau ILB est effectuée automatiquement dans le cadre de l'installation d'Oracle Solaris. Cependant, vous devez effectuer l'installation de l'utilisateur d'ILB à l'aide de la commande suivante :

```
# pkg install ilb
```

Configuration d'ILB à l'aide de la CLI

L'interface de ligne de commande se trouve dans le répertoire `/usr/sbin/ilbadm`. Elle inclut des sous-commandes pour configurer des règles d'équilibrage de la charge, des groupes de serveurs et des vérifications de l'état. D'autres sous-commandes vous permettent de consulter les statistiques et les détails de configuration. Vous devez configurer une autorisation utilisateur pour la sous-commandes de configuration ILB, à l'exception des sous-commandes d'affichage, comme `ilbadm show-rule`, `ilbadm show-server` et `ilbadm show-healthcheck`. Vous devez posséder l'autorisation RBAC `solaris.network.ilb.config` pour exécuter les sous-commandes de configuration ILB.

- Pour obtenir plus d'informations sur l'octroi de cette autorisation à un utilisateur, reportez-vous [Chapitre 3, “Affectation de droits dans Oracle Solaris”](#) du manuel “Sécurisation des utilisateurs et des processus dans Oracle Solaris 11.2”.
- Vous pouvez également attribuer cette autorisation lors de la création d'un nouveau compte utilisateur sur le système.

La commande ci-dessous crée l'utilisateur `ilbadm` qui fait partie du groupe `10`, portant l'ID `1210` et autorisé à administrer l'équilibreur de charge intégré sur le système.

```
# useradd -g 10 -u 1210 -A solaris.network.ilb.config ilbadm
```

La commande `useradd` ajoute un nouvel utilisateur aux fichiers `/etc/passwd`, `/etc/shadow` et `/etc/user_attr`. L'option `-A` affecte l'autorisation à l'utilisateur.

Les sous-commandes sont classées en deux catégories :

- **Sous-commandes de configuration** : ces sous-commandes permettant d'effectuer les tâches suivantes :
 - Création et suppression de règles d'équilibrage de la charge
 - Activation et désactivation de règles d'équilibrage de la charge
 - Création et suppression de groupes de serveurs
 - Ajout et retrait de serveurs dans un groupe de serveurs
 - Activation et désactivation de serveurs backend
 - Création et suppression de vérifications de l'état d'un groupe de serveurs dans une règle d'équilibrage de la charge

Remarque - Vous devez disposer des privilèges nécessaires pour administrer les sous-commandes de configuration. Pour créer le rôle adéquat et l'attribuer à un utilisateur, reportez-vous à [“ Création d'un rôle ” du manuel “ Sécurisation des utilisateurs et des processus dans Oracle Solaris 11.2 ”](#).

- **sous-commandes d'affichage** : ces sous-commandes vous permettent d'effectuer les tâches suivantes :
 - Affichage des règles d'équilibrage de la charge, des groupes de serveurs et des vérifications de l'état
 - Affichage des statistiques de transfert de paquets
 - Affichage de la table de connexions NAT
 - Affichage des résultats des vérifications de l'état
 - Affichage de la table de mise en correspondance de la persistance de session
-

Remarque - Aucun privilège n'est requis pour administrer les sous-commandes d'affichage.

Pour plus d'informations sur les sous-commandes `ilbadm`, reportez-vous à la page du manuel [ilbadm\(1M\)](#).

Activation ou désactivation d'ILB

Cette section indique comment activer ILB après son installation ou le désactiver si les services ILB ne sont pas obligatoires.

▼ Activation de l'équilibreur de charge intégré

1. Connexion en tant qu'administrateur.

Pour plus d'informations, reportez-vous à la section [“ A l'aide de vos droits administratifs attribués ” du manuel “ Sécurisation des utilisateurs et des processus dans Oracle Solaris 11.2 ”](#).

2. Activez le service de transfert adéquat (IPv4 ou IPv6, ou les deux).

Cette commande ne génère pas de sortie quand son exécution est réussie.

```
# ipadm set-prop -p forwarding=on ipv4
```

```
# ipadm set-prop -p forwarding=on ipv6
```

3. Activez le service ILB.

```
# svcadm enable ilb
```

4. Vérifiez que le service ILB est activé.

```
# svcs ilb
```

Cette commande affiche des informations sur le nombre d'instances du service tel qu'enregistré dans le référentiel de configuration de service.

▼ Désactivation de l'équilibreur de charge intégré

Lorsque les services d'ILB ne sont pas obligatoires, vous pouvez les désactiver.

1. Connexion en tant qu'administrateur.

Pour plus d'informations, reportez-vous à la section “ [A l'aide de vos droits administratifs attribués](#) ” du manuel “ [Sécurisation des utilisateurs et des processus dans Oracle Solaris 11.2](#) ”.

2. Désactivez le service ILB.

```
# svcadm disable ilb
```

3. Vérifiez que le service ILB est désactivé.

```
# svcs ilb
```

Cette commande affiche des informations sur le nombre d'instances du service tel qu'enregistré dans le référentiel de configuration de service.

Étapes suivantes Une fois le service ILB est désactivé, vous devez désactiver le transfert IP s'il n'est pas obligatoire.

Gestion d'un ILB

Vous pouvez configurer un ILB en définissant les groupes de serveurs, en consultant les vérifications d'intégrité d'ILB et en créant des règles ILB après avoir activé l'ILB.

Cette section traite les sujets suivants :

- “Définition des groupes de serveurs et des serveurs backend dans ILB” à la page 61
- “Surveillance de l'état de fonctionnement dans ILB” à la page 64
- “Configuration des règles ILB” à la page 67

Définition des groupes de serveurs et des serveurs backend dans ILB

Cette section explique comment créer un groupe de serveurs ILB et ajouter des serveurs backend au groupe de serveurs. Lorsqu'un serveur est ajouté à l'aide de ces sous-commandes, les ID de serveur `create-servergroup` ou `add-server` sont générés par le système. Les ID de serveurs sont uniques dans le groupe de serveurs. Pour plus d'informations sur les sous-commandes `ilbadm`, reportez-vous à la page du manuel [ilbadm\(1M\)](#).

Création d'un groupe de serveurs ILB

Pour créer un groupe de serveurs ILB, identifiez d'abord les serveurs à inclure dans le groupe de serveurs. Vous pouvez spécifier les serveurs avec leur nom d'hôte ou adresse IP et un port optionnel. Puis en tant qu'administrateur, exécutez la commande suivante :

```
# ilbadm create-servergroup -s servers=server1,server2,server3 servergroup
```

Des ID de serveur uniques préfixés par le caractère de soulignement (`_`) sont générés pour chaque serveur ajouté.

Remarque - Un serveur peut porter plusieurs ID s'il appartient à différents groupes de serveurs.

Ajout de serveurs backend à un groupe de serveurs ILB

Pour ajouter un serveur d'arrière-plan à un groupe de serveurs, connectez-vous en tant qu'administrateur et exécutez la commande suivante :

```
# ilbadm add-server -s server=server1[,server2...] servergroup
```

Les spécifications du serveur doivent inclure un nom d'hôte ou une adresse IP et peuvent éventuellement comprendre un port ou une plage de ports. Les entrées portant la même adresse IP ne sont pas autorisées au sein d'un groupe de serveurs. Des ID de serveur uniques préfixés par le caractère de soulignement (`_`) sont générés pour chaque serveur ajouté.

Remarque - Il faut indiquer les adresses IPv6 entre crochets.

EXEMPLE 6-1 Création d'un groupe de serveurs ILB et ajout de serveurs backend

L'exemple de commande ci-dessous crée un groupe nommé `webgroup` composé de trois serveurs backend.

```
# ilbadm create-servergroup -s \
servers=192.168.89.11,192.168.89.12,192.168.89.13 webgroup
# ilbadm show-servergroup
```

SGNAME	SERVERID	MINPORT	MAXPORT	IP_ADDRESS
webgroup	_webgroup.0	--	--	192.168.89.11
webgroup	_webgroup.1	--	--	192.168.89.12
webgroup	_webgroup.2	--	--	192.168.89.13

L'exemple suivant crée un groupe de serveurs appelé `webgroup1` et ajoute trois serveurs backend au groupe de serveurs.

```
# ilbadm create-servergroup webgroup1
# ilbadm add-server -s server=[2001:0db8:7::feed:6]:8080,\
[2001:0db8:7::feed:7]:8080,[2001:0db8:7::feed:8]:8080 webgroup1
```

Activation ou désactivation d'un serveur backend au sein d'un groupe de serveurs ILB

Déterminez l'adresse IP, le nom d'hôte ou l'ID du serveur à réactiver ou désactiver. Vous devez associer le groupe de serveurs avec une règle avant les serveurs du groupe de serveurs peut être activée ou désactivée.

Un serveur peut porter plusieurs ID s'il appartient à différents groupes de serveurs. Vous devez indiquer un ID de serveur pour réactiver ou désactiver le serveur pour les règles spécifiques associées à cet ID.

- Pour désactiver un serveur activé, exécutez la commande suivante :

```
# ilbadm disable-server server1
```

Le serveur sélectionné, qui est activé, est désactivé. Le noyau ne transfère pas de trafic vers ce serveur.

- Pour réactiver le serveur désactivé, exécutez la commande suivante :

```
# ilbadm enable-server server1
```

Le serveur sélectionné, qui est désactivé, est réactivé.

- Pour afficher l'état du serveur, exécutez la commande suivante :

```
# ilbadm show-server [[-p] -o field[,field...]] [rulename]
```

Remarque - Un serveur affiche l'état comme étant activé ou désactivé uniquement lorsque le groupe de serveurs auquel le serveur appartient est associé à la règle.

EXEMPLE 6-2 Réactivation ou désactivation d'un serveur backend au sein d'un groupe de serveurs ILB

Dans l'exemple ci-dessous, un serveur portant l'ID `websg.1` est désactivé, puis réactivé.

```
# ilbadm enable-server _websg.1
# ilbadm disable-server _websg.1
```

Suppression d'un serveur backend d'un groupe de serveurs ILB

Vous pouvez supprimer un serveur d'arrière-plan d'un groupe de serveurs ILB ou de tous les groupes de serveurs à l'aide de la commande `ilbadm remove-server`. Déterminez l'ID du serveur à supprimer du groupe de serveurs.

```
ilbadm show-servergroup -o all
```

L'ID de serveur est un nom unique associé à l'adresse IP attribuée à un système au moment de l'ajout du serveur à un groupe de serveurs.

Ensuite, supprimez le serveur.

```
# ilbadm remove-server -s server=server-ID server-group
```

Si une règle NAT ou Half-NAT concerne le serveur, désactivez-le en exécutant la sous-commande `disable-server` avant de le supprimer. Pour plus d'informations, reportez-vous à la section [“Activation ou désactivation d'un serveur backend au sein d'un groupe de serveurs ILB” à la page 62](#). Lorsqu'il est désactivé, un serveur passe à l'état `connection-draining`. Vérifiez périodiquement la table NAT à l'aide de la commande `ilbadm show-nat` pour vérifier si le serveur possède toujours des connexions. Une fois que toutes les connexions ont été vidées (le serveur n'est pas affiché dans la sortie de commande `show-nat`), vous pouvez supprimer le serveur à l'aide de la commande `remove-server`.

Si la valeur `conn-drain` est définie, l'état `connection-draining` prend fin au terme de ce délai d'expiration. La valeur par défaut du délai `conn-drain` est fixée à 0, ce qui signifie que le serveur reste en attente jusqu'à ce qu'une connexion soit progressivement interrompue.

EXEMPLE 6-3 Suppression d'un serveur backend d'un groupe de serveurs ILB

L'exemple de commande ci-dessous supprime le serveur portant l'ID `_sg1.2` du groupe `sg1`.

```
# ilbadm remove-server -s server=_sg1.2 sg1
```

Suppression de groupes de serveurs ILB

Cette section décrit la procédure à suivre pour supprimer un groupe de serveurs ILB. Vous ne pouvez pas supprimer un groupe de serveurs qui est utilisé par n'importe quelle règle active.

Tout d'abord, affichez les informations disponibles sur tous les groupes de serveurs.

```
# ilbadm show-servergroup -o all
sgname      serverID      minport      maxport      IP_address
specgroup   _specgroup.0  7001         7001         192.168.68.18
specgroup   _specgroup.1  7001         7001         192.168.68.19
test123     _test123.0    7002         7002         192.168.67.18
test123     _test123.1    7002         7002         192.168.67.19
```

Exécutez la commande suivante :

```
# ilbadm delete-servergroup servergroup
```

Si le groupe de serveurs est concerné par une règle active, la suppression échoue.

L'exemple de commande ci-dessous supprime le groupe nommé `webgroup`.

```
# ilbadm delete-servergroup webgroup
```

Surveillance de l'état de fonctionnement dans ILB

ILB fournit les types de contrôle de l'intégrité du serveur facultatifs suivants :

- Tests ping intégrés
- Tests TCP intégrés
- Tests UDP intégrés
- Tests personnalisés fournis par l'utilisateur pouvant être exécutés comme des vérifications de l'état

Par défaut, ILB n'effectue aucune vérification de l'état. Vous pouvez spécifier des vérifications de l'état de chaque groupe de serveurs lors de la création d'une règle d'équilibrage de la charge. Vous pouvez configurer une seule vérification de l'état par règle d'équilibrage de la charge.

Tant qu'un service virtuel est activé, les contrôles de l'intégrité définis sur le groupe de serveurs associé à ce service virtuel démarrent automatiquement et se répètent de manière périodique. Les vérifications de l'état cessent dès que le service virtuel est désactivé. Notez que les statuts de vérification de l'état ne sont pas conservés à la réactivation du service virtuel.

Lorsque vous spécifiez un test TCP, UDP ou personnalisé pour vérifier l'état d'un serveur, l'équilibreur de charge intégré envoie une commande ping par défaut pour déterminer s'il est accessible avant d'envoyer le test proprement dit. En cas d'échec du test ping, le serveur est désactivé, et son état passe à `unreachable`. Si la commande ping aboutit, mais que le test TCP, UDP ou personnalisé échoue, le serveur est désactivé et son état passe à `dead`.

Vous pouvez désactiver le test ping par défaut UDP à l'exception des du test UDP. Le test ping est toujours le test par défaut pour les vérifications d'état UDP.

Création d'une vérification de l'état

Vous pouvez créer une vérification de l'état et affecter la vérification de l'état à un groupe de serveurs lors de la création d'une règle d'équilibrage de charge. L'exemple de commande ci-dessous montre comment sont créés les objets de vérification de l'état `hc1` et `hc-myscript`. La première vérification de l'état met en oeuvre le test TCP intégré, Le second contrôle de l'intégrité utilise un test personnalisé : `/var/tmp/my-script`.

```
# ilbadm create-healthcheck -h hc-timeout=3,\
hc-count=2,hc-interval=8,hc-test=tcp hc1
# ilbadm create-healthcheck -h hc-timeout=3,\
hc-count=2,hc-interval=8,hc-test=/var/tmp/my-script hc-myscript
```

Les arguments sont les suivants :

<code>hc-timeout</code>	Spécifie le délai d'expiration au terme duquel le système considère que la vérification de l'état a échoué si elle n'est pas terminée.
<code>hc-count</code>	Spécifie le nombre de tentatives d'exécution de la vérification de l'état <code>hc-test</code> .
<code>hc-interval</code>	Spécifie l'intervalle de temps à respecter entre deux vérifications consécutives. Pour éviter l'envoi de sondes à tous les serveurs en même temps, l'intervalle est défini de manière aléatoire entre $0.5 * hc-interval$ et $1.5 * hc-interval$.
<code>hc-test</code>	Indique le type de vérification de l'état. Vous pouvez spécifier les vérifications d'état intégrées, telles que <code>tcp</code> , <code>udp</code> et <code>ping</code> ou une vérification d'état externe qui doit être indiquée avec son chemin d'accès complet.

Remarque - Le port réservé au test `hc-test` est spécifié avec le mot clé `hc-port` dans la sous-commande `create-rule`. Pour plus d'informations, reportez-vous à la page de manuel [ilbadm\(1M\)](#).

Un test personnalisé fourni par l'utilisateur peut être un fichier binaire ou un script.

- Le test peut résider n'importe où sur le système. Vous devez indiquer le chemin d'accès absolu par le biais de la sous-commande `create-healthcheck`.

Lorsque vous définissez ce test (par exemple `/var/tmp/my-script`) pour le contrôle de l'intégrité dans la sous-commande `create-rule`, le démon `ilbd` clone un processus et exécute le test comme suit :

```
/var/tmp/my-script $1 $2 $3 $4 $5
```

Les arguments sont les suivants :

- | | |
|-----|--|
| \$1 | VIP (adresse IPv4 ou IPv6 littérale) |
| \$2 | IP du serveur (adresse IPv4 ou IPv6 littérale) |
| \$3 | Protocole (UDP, TCP sous forme de chaîne) |
| \$4 | Plage de ports numérique (valeur spécifiée par l'utilisateur pour <code>hc-port</code>) |
| \$5 | Durée maximale (en secondes) à respecter avant de considérer que le test a échoué. Si le test n'est pas terminé au terme de ce délai, il peut être interrompu et associé à un statut d'échec. Cette valeur définie par l'utilisateur est indiquée dans <code>hc-timeout</code> . |
- Le test fourni par l'utilisateur peut comprendre la totalité ou une partie des arguments, mais il doit *impérativement* renvoyer une des informations suivantes :
 - Temps d'aller-retour réseau (en microsecondes)
 - 0 si le test ne calcule pas le temps d'aller-retour réseau
 - -1 en cas d'échec

Par défaut, le test de vérification de l'état s'exécute avec les privilèges: `PRIV_PROC_FORK`, `RIV_PROC_EXEC` et `RIV_NET_ICMPACCESS`.

Si un ensemble de privilèges plus large est requis, il faut implémenter `setuid` dans le test. Pour plus d'informations sur les privilèges, reportez-vous à la page de manuel [privileges\(5\)](#).

Liste des contrôles de l'intégrité

Pour obtenir des informations détaillées sur les vérifications configurées, exécutez la commande suivante :

```
# ilbadm show-healthcheck
HCNAME    TIMEOUT COUNT  INTERVAL DEF_PING TEST
hc1       3       2      8        Y      tcp
hc2       3       2      8        N      /var/usr-script
```

Affichage des résultats des vérifications de l'état

Utilisez la commande `ilbadm list-hc-result` pour obtenir les résultats des contrôles de l'intégrité. Si vous ne spécifiez pas de règle ou de vérification de l'état, la sous-commande répertorie toutes les vérifications configurées.

L'exemple de commande ci-dessous affiche les résultats des vérifications de l'état associées à une règle nommée `rule1`.

```
# ilbadm show-hc-result rule1
RULENAME  HCNAME  SERVERID  STATUS  FAIL LAST      NEXT      RTT
rule1     hc1     _sg1:0    dead    10  11:01:19  11:01:27  941
rule1     hc1     _sg1:1    alive   0   11:01:20  11:01:34  1111
```

Remarque - La sous-commande `show-hc-result` affiche la vérification de l'état uniquement quand les règles sont associées à des vérifications de l'état.

La colonne `LAST` du tableau indique l'heure de la dernière vérification de l'état du serveur. La colonne `NEXT` affiche l'heure à laquelle la prochaine vérification de l'état sera effectuée.

Suppression d'une vérification de l'état

Vous pouvez supprimer une vérification de l'état à l'aide de la commande `ilbadm delete-healthcheck`. L'exemple de commande ci-dessous supprime une vérification de l'état nommée `hc1` :

```
# ilbadm delete-healthcheck hc1
```

Configuration des règles ILB

Cette section indique comment exécuter la commande `ilbadm` pour créer, supprimer et répertorier les règles d'équilibrage de la charge.

Algorithmes ILB

Les algorithmes ILB contrôlent la distribution du trafic et fournissent diverses caractéristiques pour répartir la charge et sélectionner les serveurs.

L'équilibreur de charge intégré fournit les algorithmes suivants dans les deux topologies :

- Tourniquet : avec cet algorithme (round-robin), l'équilibreur de charge affecte les demandes à un groupe de serveurs à tour de rôle. Après l'attribution d'une demande à un serveur, celui-ci est placé au bas de la liste.
- Hachage *src IP* : avec la méthode de hachage de l'adresse IP source, l'équilibreur de charge sélectionne un serveur en fonction de la valeur de hachage associée à l'adresse IP source de la demande entrante.
- Hachage *src-IP, port* : avec la méthode de hachage du port et de l'adresse IP source, l'équilibreur de charge sélectionne un serveur en fonction de la valeur de hachage associée au port et à l'adresse IP source de la demande entrante.
- Hachage *src-IP, VIP* : avec la méthode de hachage de l'adresse IP source et VIP, l'équilibreur de charge sélectionne un serveur en fonction de la valeur de hachage associée à l'adresse IP source et de destination de la demande entrante.

Création d'une règle de ILB

Dans ILB, un service virtuel est représenté par une règle d'équilibrage de charge et défini par les paramètres suivants :

- Adresse IP virtuelle
- Protocole de transport (TCP ou UDP)
- Numéro de port (ou plage de ports)
- Algorithme d'équilibrage de charge
- Mode d'équilibrage de charge (DSR, Full-NAT ou Half-NAT)
- Groupe composé de plusieurs serveurs backend
- Vérifications de l'état (en option) à exécuter sur chacun des serveurs du groupe
- Port (facultatif) sur lequel réaliser les vérifications de l'état

Remarque - Vous pouvez spécifier des contrôles de l'intégrité sur un port particulier, ou sur n'importe quel port sélectionné de manière aléatoire par le démon `ilbd` dans la plage de numéros de port associés au serveur.

- Nom de règle pour représenter un service virtuel

Avant de pouvoir créer une règle, vous devez procéder comme suit :

- Créez un groupe qui inclut les serveurs backend appropriés. Pour plus d'informations, reportez-vous à la section “[Définition des groupes de serveurs et des serveurs backend dans ILB](#)” à la page 61.
- Créez une vérification de l'état afin d'associer le serveur de données avec la règle de vérification de l'état. Pour plus d'informations, reportez-vous à la section “[Création d'une vérification de l'état](#)” à la page 65.
- Déterminez l'adresse VIP, le port et le protocole en option à associer à la règle.
- Sélectionnez l'opération à mettre en place (DSR, Half-NAT ou Full-NAT).
- Identifiez l'algorithme d'équilibrage de charge à utiliser. Pour plus d'informations, reportez-vous à la section “[Algorithmes ILB](#)” à la page 68.

Vous pouvez créer une règle ILB à l'aide de la commande `ilbadm create-rule`. Pour plus d'informations sur l'utilisation de la commande `ilbadm create-rule`, reportez-vous à la page de manuel [ilbadm\(1M\)](#).

Respectez la syntaxe suivante :

```
# ilbadm create-rule -e -i vip=IPAddr,port=port,protocol=protocol \
-m lbalg=lb-algorithm,type=topology-type,proxy-src=IPAddr1-IPAddr2,\
pmask=value -h hc-name=hc1-o servergroup=sg rule1
```

Remarque - L'option `-e` active la règle qui est en cours de création, et qui, autrement, serait désactivée par défaut.

EXEMPLE 6-4 Création d'une règle Full-NAT avec persistance de la session de vérification de l'état

Cet exemple crée une vérification de l'état nommée `hc1` et un groupe de serveurs nommé `sg1`. Le groupe se compose de deux serveurs, chacun disposant d'une plage de ports. La dernière commande crée et active une règle nommée `rule1`, puis l'associe au groupe de serveurs et à la vérification de l'état. Cette règle implémente le mode Full-NAT. Notez que la création du groupe de serveurs et de la vérification de l'état doit précéder la création de la règle.

```
# ilbadm create-healthcheck -h hc-test=tcp,hc-timeout=2,\
hc-count=3,hc-interval=10 hc1
# ilbadm create-servergroup -s server=192.168.0.10:6000-6009,192.168.0.11:7000-7009 sg1
# ilbadm create-rule -e -p -i vip=10.0.0.10,port=5000-5009,\
protocol=tcp -m lbalg=rr,type=NAT,proxy-src=192.168.0.101-192.168.0.104,pmask=24 \
-h hc-name=hc1 -o servergroup=sg1 rule1
```

Après avoir créé une correspondance persistante, les demandes de connexions et/ou de paquets suivantes sur un service virtuel, avec une adresse IP source du client correspondante, sont transférées au même serveur d'arrière-plan. La longueur du préfixe de la notation CIDR (Classless Inter-Domain Routing) correspond à une valeur comprise dans la plage 0-32 pour IPv4 ou 0-128 pour IPv6.

Lorsque vous créez une règle Half-NAT ou Full-NAT, spécifiez la valeur du délai d'expiration connection-drain. La valeur par défaut du délai conn-drain est fixée à 0, ce qui signifie que le serveur reste en attente jusqu'à ce qu'une connexion soit progressivement interrompue.

Etablissement de la liste des règles ILB

Pour répertorier les détails de configuration d'une règle, émettez la commande suivante :
Si vous ne spécifiez pas de nom de règle, les informations relatives à toutes les règles sont renvoyées.

```
# ilbadm show-rule
```

RULENAME	STATUS	LBALG	TYPE	PROTOCOL	VIP	PORT
rule-http	E	hash-ip-port	NAT	TCP	10.0.0.1	80
rule-dns	D	hash-ip	NAT	UDP	10.0.0.1	53
rule-abc	D	roundrobin	NAT	TCP	2001:db8::1	1024
rule-xyz	E	ip-vip	NAT	TCP	2001:db8::1	2048-2050

Suppression d'une règle ILB

Vous pouvez utiliser la commande `ilbadm delete-rule` pour supprimer une règle. Ajoutez l'option `-a` pour supprimer l'intégralité des règles. L'exemple de commande ci-dessous supprime la règle nommée `rule1`.

```
# ilbadm delete-rule rule1
```

Cas d'utilisation : configuration d'un ILB

Cette section décrit les étapes à suivre pour configurer l'équilibreur de charge intégré de manière à utiliser une topologie Half-NAT pour répartir le trafic sur deux serveurs. Reportez-vous à l'implémentation de la topologie NAT à la section [“Modes de fonctionnement d'ILB” à la page 50](#).

1. Connexion en tant qu'administrateur.

Pour plus d'informations, reportez-vous à la section [“ A l'aide de vos droits administratifs attribués ” du manuel “ Sécurisation des utilisateurs et des processus dans Oracle Solaris 11.2 ”](#).

2. Configurez le groupe de serveurs dans ILB.

Les deux serveurs sont 192.168.1.50 et 192.169.1.60. Vous pouvez créer le groupe de serveurs `srvgrp1` constitué de ces deux serveurs en exécutant la commande suivante. Pour plus d'informations sur la configuration d'un groupe de serveurs dans ILB, reportez-vous à la section [“Création d'un groupe de serveurs ILB” à la page 61](#).

```
# ilbadm create-sg -s servers=192.168.1.50,192.168.1.60 srvgrp1
```

3. Configurez les serveurs backend.

Les serveurs backend sont configurés de façon à utiliser ILB en tant que routeur par défaut dans ce scénario. Exécutez la commande suivante sur les deux serveurs, procédez comme suit :

```
# route add -p default 192.168.1.21
```

Après avoir exécuté cette commande, démarrez les applications serveur sur les deux serveurs. Supposons qu'il s'agit d'une application TCP en mode écoute sur le port 5000. Pour plus d'informations sur la configuration des serveurs backend, reportez-vous à la section [“Ajout de serveurs backend à un groupe de serveurs ILB” à la page 61.](#)

4. Configurez une vérification de l'état simple nommée hc-srvgrp1. Créez la vérification de l'état en exécutant la commande suivante :

```
# ilbadm create-hc -h hc-test=tcp,hc-timeout=3,\
hc-count=3,hc-interval=60 hc-srvgrp1
```

Une vérification de l'état de niveau TCP permet de déterminer si l'application de serveur est accessible. Cette vérification est effectuée toutes les 60 secondes. L'opération effectuée au plus 3 tentatives et respecte un intervalle de trois secondes au maximum entre chaque tentative pour déterminer si un serveur est en fonctionnel. Si les trois tentatives échouent, comme périmée il marque le serveur dead. Pour plus d'informations sur la surveillance et la création de vérifications d'état, reportez-vous à la section [“Surveillance de l'état de fonctionnement dans ILB” à la page 64.](#)

5. Configurez une règle ILB en exécutant la commande ci-dessous :

```
# ilbadm create-rule -e -p -i vip=10.0.2.20,port=5000 -m \
lbalg=rr,type=half-nat,pmask=32 \
-h hc-name=hc-srvgrp1 -o servergroup=srvgrp1 rule1_rr
```

La persistance (avec un masque à 32 bits) est appliquée dans cette règle. L'algorithme de l'équilibreur de charge est round robin. Pour plus d'informations sur les différents algorithmes ILB, reportez-vous à [“Algorithmes ILB” à la page 68.](#) Le groupe de serveurs concerné est srvgrp1 et le mécanisme de vérification de l'état mis en oeuvre est hc-srvgrp1. Pour plus d'informations sur la création de règles ILB, reportez-vous à la section [“Création d'une règle de ILB” à la page 68.](#)

Affichage des statistiques ILB

Cette section indique comment exécuter la commande `ilbadm` pour obtenir des informations telles que les statistiques d'impression d'un serveur ou les statistiques relatives à une règle. Vous

pouvez également consulter la table de connexions NAT et la table de mise en correspondance de la persistance de session.

Affichage des données statistiques

Utilisez la commande `ilbadm show-statistics` pour visualiser les informations sur la distribution de charge comme indiqué dans l'exemple suivant.

```
# ilbadm show-statistics
PKT_P  BYTES_P  PKT_U  BYTES_U  PKT_D  BYTES_D
9       636       0       0       0       0

PKT_P                Paquets traités
BYTES_P              Octets traités
PKT_U                Paquets non traités
BYTES_U              Octets non traités
PKT_D                Paquets rejetés
BYTES_D              Octets rejetés
```

Affichage de la table de connexions NAT

Exécutez la commande `ilbadm show-nat` pour afficher la table de connexions NAT. Notez que les positions relatives des éléments et les exécutions consécutives de cette commande ne sont pas prises en compte. Par exemple, si vous lancez la commande `ilbadm show-nat 10` deux fois, vous risquez de ne pas voir les mêmes 10 éléments chaque fois que vous exécutez, surtout si le système est occupé. Si vous ne spécifiez pas de nombre, la table de connexion NAT s'affiche dans son intégralité.

EXEMPLE 6-5 Entrées de la table de connexions NAT

L'exemple ci-dessous répertorie cinq entrées de la table de connexions NAT.

```
# ilbadm show-nat 5
UDP: 124.106.235.150.53688 > 85.0.0.1.1024 >>> 82.0.0.39.4127 > 82.0.0.56.1024
UDP: 71.159.95.31.61528 > 85.0.0.1.1024 >>> 82.0.0.39.4146 > 82.0.0.55.1024
UDP: 9.213.106.54.19787 > 85.0.0.1.1024 >>> 82.0.0.40.4114 > 82.0.0.55.1024
UDP: 118.148.25.17.26676 > 85.0.0.1.1024 >>> 82.0.0.40.4112 > 82.0.0.56.1024
```

UDP: 69.219.132.153.56132 > 85.0.0.1.1024 >>> 82.0.0.39.4134 > 82.0.0.55.1024

Les entrées respectent le format suivant :

T: IP1 > IP2 >>> IP3 > IP4

T	Protocole de transport utilisé dans cette entrée
IP1	Adresse IP et port du client
IP2	VIP et port
IP3	En mode Half-NAT, adresse IP et port du client En mode Full-NAT, adresse IP et port du client
IP4	Adresse IP et port du serveur backend

Affichage de la table de mise en correspondance de la persistance de session

Exécutez la sous-commande `ilbadm show-persist` pour afficher la table de mise en correspondance de la persistance de session.

EXEMPLE 6-6 Entrées de la table de mise en correspondance de la persistance de session

L'exemple ci-dessous répertorie cinq entrées de la table de mise en correspondance de la persistance de session.

```
# ilbadm show-persist 5
rule2: 124.106.235.150 --> 82.0.0.56
rule3: 71.159.95.31 --> 82.0.0.55
rule3: 9.213.106.54 --> 82.0.0.55
rule1: 118.148.25.17 --> 82.0.0.56
rule2: 69.219.132.153 --> 82.0.0.55
```

Les entrées respectent le format suivant :

R: IP1 --> IP2

R	Règle à laquelle est associée cette entrée de persistance
IP1	Adresse IP du client
IP2	Adresse IP du serveur backend

Importation et exportation des configurations

Les sous-commandes `export` et `import` sont utilisées pour transférer une configuration d'un système vers un autre. Par exemple, si vous souhaitez configurer une sauvegarde d'ILB à l'aide de VRRP afin qu'il dispose d'une configuration active-passive, vous pouvez vous contenter d'exporter la configuration dans un fichier et l'importer dans le système de sauvegarde. La commande `ilbadm export` exporte la configuration ILB actuelle dans un fichier spécifié par l'utilisateur. Vous pouvez ensuite utiliser ces informations en tant qu'entrée de la commande `ilbadm import`.

Sauf instruction expresse imposant sa conservation, la commande `ilbadm import` supprime la configuration existante avant l'importation. Si vous ne spécifiez pas de nom de fichier, la commande lit le flux d'entrée `stdin` ou écrit dans le flux de sortie `stdout`.

Pour exporter une configuration ILB, exécutez la commande `export-config`. L'exemple ci-dessous exporte la configuration actuelle dans le fichier `/var/tmp/ilb_config` dans un format adapté à l'importation à l'aide de la commande `import` :

```
# ilbadm export-config /var/tmp/ilb_config
```

Pour importer une configuration ILB, exécutez la commande `import-config`. L'exemple suivant lit le contenu du fichier `/var/tmp/ilb_config` et remplace la configuration existante :

```
# ilbadm import-config /var/tmp/ilb_config
```

Configuration ILB pour une haute disponibilité

Ce chapitre décrit la configuration de haute disponibilité (HA) d'ILB à l'aide de la fonction VRRP. ILB est configuré pour la haute disponibilité à l'aide des topologies DSR et Half-NAT. Les topologies DSR et Half-NAT utilisent VRRP afin de protéger l'adresse IP virtuelle d'une règle ILB. Toutefois, dans la topologie Half-NAT, VRRP est également utilisé pour protéger l'adresse IP de la base de données de l'équilibreur de charge principal connecté aux serveurs backend. Cela permet de vous assurer que lorsque l'équilibreur de charge principal échoue, les serveurs backend utilisent l'équilibreur de charge de secours (passif).

Pour plus d'informations sur VRRP, voir [Chapitre 3, Utilisation de Virtual Router Redundancy Protocol](#) et pour plus d'informations sur la configuration et la gestion d'ILB, voir [Chapitre 6, Configuration et gestion de l'équilibreur de charge intégré](#).

Ce chapitre aborde les sujets suivants :

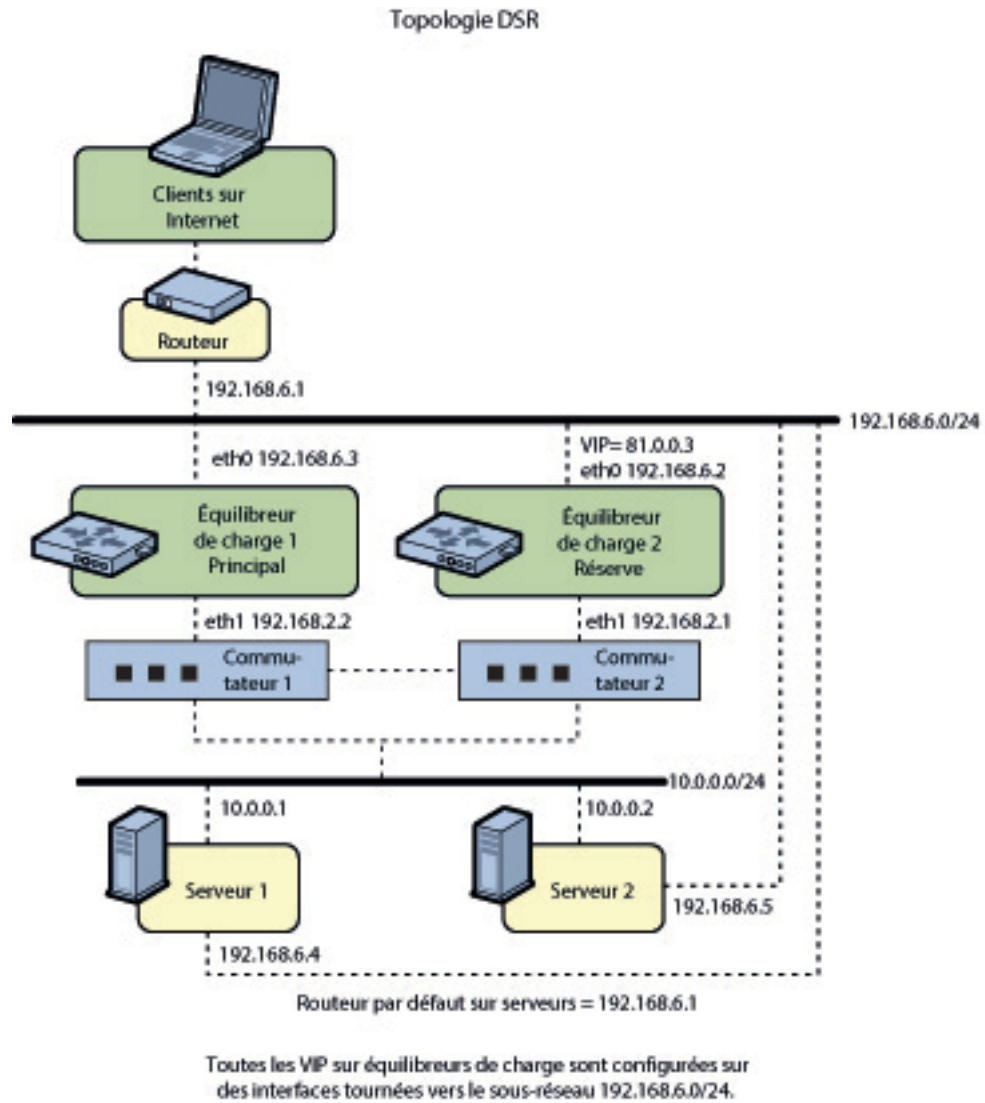
- “Configuration d'ILB pour la haute disponibilité à l'aide de la topologie DSR” à la page 75
- “Configuration d'ILB pour la haute disponibilité à l'aide de la topologie Half-NAT” à la page 78

Configuration d'ILB pour la haute disponibilité à l'aide de la topologie DSR

Il faut configurer deux équilibreurs de charge, le premier faisant office d'équilibreur principal et le deuxième d'équilibreur de secours. Le programme d'équilibrage de charge principal agit en tant que routeur maître et le programme d'équilibrage de charge de secours (passif) agit en tant que routeur de secours. L'adresse IP virtuelle d'une règle ILB agit comme une adresse IP du routeur virtuel. Le sous-système VRRP vérifie si le programme d'équilibrage de charge principal a échoué. En cas de panne de l'équilibreur principal, l'équilibreur de secours prend le relais.

Le schéma suivant illustre la topologie DSR pour configurer les connexions ILB permettant d'activer la haute disponibilité.

FIGURE 7-1 Configuration d'ILB pour la haute disponibilité à l'aide de la topologie DSR



▼ Configuration d'ILB pour la haute disponibilité à l'aide de la topologie DSR

Vous pouvez configurer les programmes d'équilibrage de charge principal et de secours de la même façon pour la règle ILB, le groupe de serveurs et la vérification de l'état. Vous pouvez configurer ces deux équilibreurs de charge pour utiliser le VRRP. Définissez aussi l'adresse IP virtuelle de la règle comme adresse de routeur virtuel. Le sous-système VRRP s'assure ensuite que l'un des programmes d'équilibrage de charge est toujours actif.

1. Connexion en tant qu'administrateur.

Pour plus d'informations, reportez-vous à la section “ [A l'aide de vos droits administratifs attribués](#) ” du manuel “ [Sécurisation des utilisateurs et des processus dans Oracle Solaris 11.2](#) ”.

2. Configurez les équilibreurs de charge principal et d'urgence (passif) de la même manière.

```
# ilbadm create-servergroup -s server=10.0.0.1,10.0.0.2 sg1
# ilbadm create-rule -i vip=10.81.0.3,port=9001 \
-m lbalg=hash-ip-port,type=DSR -o servergroup=sg1 rule1
```

3. Configurez l'équilibreur de charge 1 en tant qu'équilibreur principal.

```
LB1# dladm create-vnic -m vrrp -V 1 -A inet -l eth0 vnic1
LB1# vrrpadm create-router -V 1 -A inet -l eth0 -p 255 vrrp1
LB1# ipadm create-ip vnic1
LB1# ipadm create-addr -d -a 10.81.0.3/24 vnic1
```

La priorité du routeur vrrp1 est définie de manière à être 255 à l'aide de la commande vrrpadm. La valeur de priorité rend le routeur maître et donc le programme d'équilibrage de charge actif.

4. Configurez l'équilibreur de charge 2 en tant qu'équilibreur de secours.

```
LB2# dladm create-vnic -m vrrp -V 1 -A inet -l eth0 vnic1
LB2# vrrpadm create-router -V 1 -A inet -l eth0 -p 100 vrrp1
LB2# ipadm create-ip vnic1
LB2# ipadm create-addr -d -a 10.81.0.3/24 vnic1
```

Cette configuration assure une protection contre les scénarios d'échec suivants :

- Si l'équilibreur de charge 1 subit une défaillance, l'équilibreur 2 prend le relais. L'équilibreur de charge 2 gère la résolution d'adresse pour le protocole VIP 10.81.0.3 et traite tous les paquets en provenance de clients à destination de l'adresse IP 10.81.0.3.
Dès que l'équilibreur de charge 1 est à nouveau opérationnel, l'équilibreur 2 repasse en mode de veille.
- En cas de panne de l'une ou des deux interfaces de l'équilibreur de charge 1, l'équilibreur 2 prend le relais. L'équilibreur de charge 2 gère ensuite la résolution d'adresse pour le

protocole VIP 10.81.0.3 et traite tous les paquets en provenance de clients à destination de l'adresse IP 10.81.0.3.

Dès que les interfaces de l'équilibreur de charge 1 sont opérationnelles, l'équilibreur 2 repasse en mode de veille.

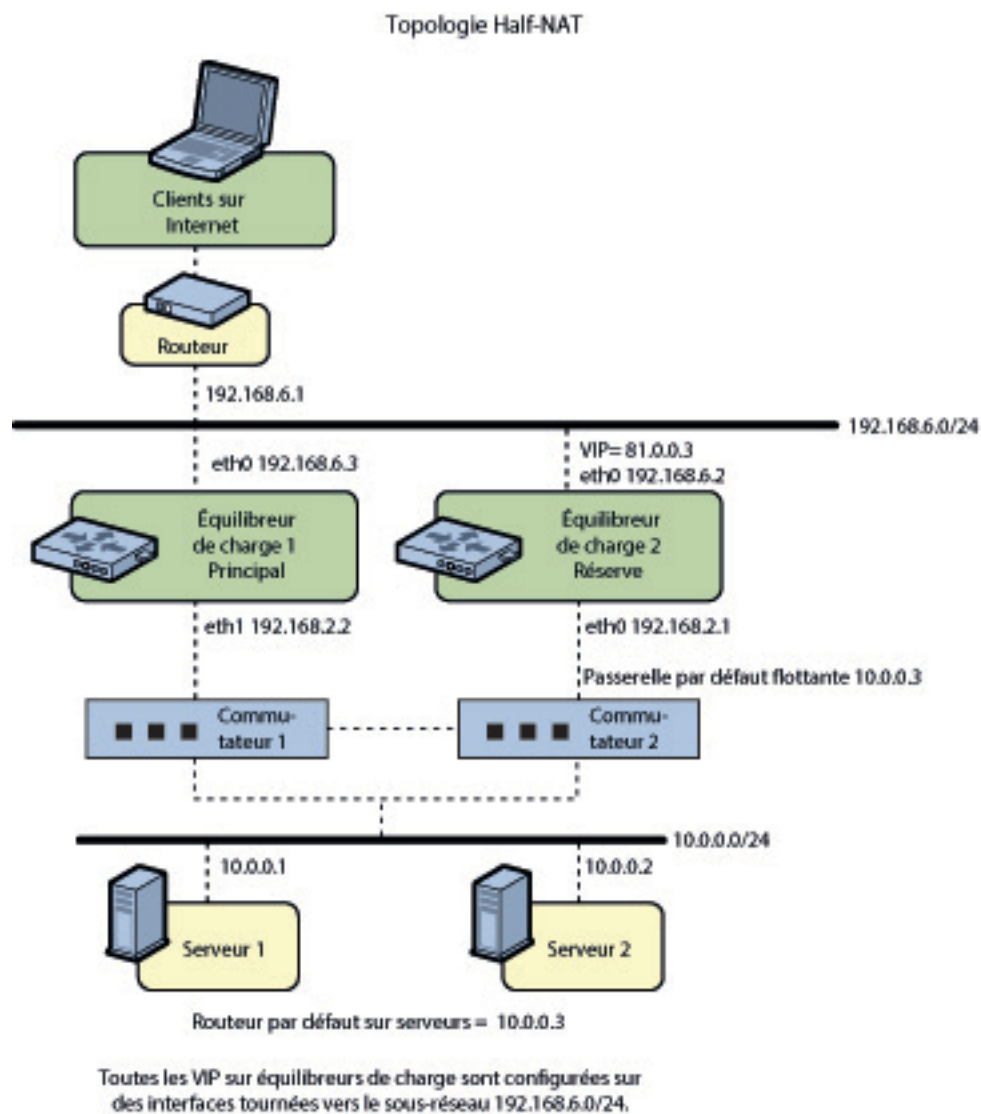
Configuration d'ILB pour la haute disponibilité à l'aide de la topologie Half-NAT

Cette section décrit la procédure de configuration des connexions ILB de manière à atteindre la haute disponibilité à l'aide de la topologie Half-NAT. Il faut configurer deux équilibreurs de charge, le premier faisant office d'équilibreur principal et le deuxième d'équilibreur de secours. En cas de panne de l'équilibreur principal, l'équilibreur de secours prend le relais.

Remarque - L'implémentation actuelle de l'équilibreur de charge intégré ne synchronise pas les données des équilibreurs principal et de secours. Lorsque l'équilibreur de charge de secours prend le relais en cas de défaillance de l'équilibreur principal, toutes les connexions en cours sont interrompues. Cela étant, la haute disponibilité sans synchronisation présente bien des avantages en cas de panne de l'équilibreur de charge principal.

Le schéma suivant illustre la topologie Half-NAT pour configurer les connexions ILB permettant d'activer la haute disponibilité.

FIGURE 7-2 Configuration de l'équilibreur de charge intégré à haute disponibilité à l'aide de la topologie Half-NAT



▼ Configuration d'ILB pour la haute disponibilité à l'aide de la topologie Half-NAT

1. Connexion en tant qu'administrateur.

Pour plus d'informations, reportez-vous à la section “ [A l'aide de vos droits administratifs attribués](#) ” du manuel “ [Sécurisation des utilisateurs et des processus dans Oracle Solaris 11.2](#) ”.

2. Configurez les équilibreurs de charge principal et de secours.

```
# ilbadm create servergroup -s server=10.0.0.1,10.0.0.2 sg1
# ilbadm create-rule -ep -i vip=10.81.0.3,port=9001-9006,protocol=udp \
-m lbalg=roundrobin,type=HALF-NAT,pmask=24 \
-h hc-name=hc1,hc-port=9006 \
-t conn-drain=70,nat-timeout=70,persist-timeout=70 -o servergroup=sg1 rule1
```

3. Configurez l'équilibreur de charge 1 en tant qu'équilibreur principal.

```
LB1# dladm create-vnic -m vrrp -V 1 -A inet -l eth0 vnic1
LB1# ipadm create-ip vnic1
LB1# ipadm create-addr -d -a 10.81.0.3/24 vnic1
LB1# vrrpadm create-router -V 1 -A inet -l eth0 -p 255 vrrp1
LB1# dladm create-vnic -m vrrp -V 2 -A inet -l eth1 vnic2
LB1# ipadm create-ip vnic2
LB1# ipadm create-addr -d -a 10.0.0.3/24 vnic2
LB1# vrrpadm create-router -V 2 -A inet -l eth1 -p 255 vrrp2
```

4. Configurez l'équilibreur de charge 2 en tant qu'équilibreur de secours.

```
LB2# dladm create-vnic -m vrrp -V 1 -A inet -l eth0 vnic1
LB2# ipadm create-ip vnic1
LB2# ipadm create-addr -d -a 10.81.0.3/24 vnic1
LB2# vrrpadm create-router -V 1 -A inet -l eth0 -p 100 vrrp1
LB2# dladm create-vnic -m vrrp -V 2 -A inet -l eth1 vnic2
LB2# ipadm create-ip vnic2
LB2# ipadm create-addr -d -a 10.0.0.3/24 vnic2
LB2# vrrpadm create-router -V 2 -A inet -l eth1 -p 100 vrrp2
```

5. Ajoutez l'adresse IP de la passerelle flottante par défaut sur les deux serveurs.

```
# route add default 10.0.0.3
```

Cette configuration assure une protection contre les scénarios d'échec suivants :

- Si l'équilibreur de charge 1 subit une défaillance, l'équilibreur 2 prend le relais. L'équilibreur de charge 2 gère la résolution d'adresse pour le protocole VIP 10.81.0.3 et traite tous les paquets en provenance de clients à destination de l'adresse IP 10.81.0.3. L'équilibreur de charge 2 traite également l'ensemble des paquets envoyés à l'adresse de la passerelle flottante 10.0.0.3.

Dès que l'équilibreur de charge 1 est à nouveau opérationnel, l'équilibreur 2 repasse en mode de veille.

- En cas de panne de l'une ou des deux interfaces de l'équilibreur de charge 1, l'équilibreur 2 prend le relais. L'équilibreur de charge 2 gère ensuite la résolution d'adresse pour le protocole VIP 10.81.0.3 et traite tous les paquets en provenance de clients à destination de l'adresse IP 10.81.0.3. L'équilibreur de charge 2 traite également l'ensemble des paquets envoyés à l'adresse de la passerelle flottante 10.0.0.3.

Dès que les interfaces de l'équilibreur de charge 1 sont opérationnelles, l'équilibreur 2 repasse en mode de veille.

Index

A

- activation
 - routeur VRRP, 41
- administration
 - ILB, 61, 64, 67
- adresses IP associées avec des routeurs VRRP
 - affichage , 44
- affichage
 - adresse IP associée avec un routeur VRRP, 44
 - configuration d'un routeur VRRP, 42
 - contrôle de l'intégrité, 67
- ajouter
 - groupe de serveurs ILB, 61

B

- BGP, 11
- bserveur arrière-plan
 - désactivation, 62

C

- client-à-serveur, 54
- comparaison de couche 2 VRRP avec couche 3 VRRP, 31
- configuration
 - adresse IP virtuelle pour un routeur VRRP, 40
 - routeurs, 10, 17
 - routeurs prenant en charge IPv6, 23
- configuration d'un routeur
 - routeur IPv4, 17
- configuration de routeur
 - routeur IPv6, 22
- configuration réseau
 - routeur, 18
 - routeur IPv6, 23

- contrôle de l'intégrité

- affichage des résultats, 67
 - création, 65
 - suppression, 67

- contrôles de l'état de fonctionnement dans ILB
 - surveillance, 64

- contrôles de l'intégrité
 - affichage, 67

- couche 2 VRRP comparée avec couche 3 VRRP, 31

- couche 3 VRRP

- limitations, 33

- présentation, 30

- Prise en charge d'Ethernet via InfiniBand, 33

- création

- contrôle de l'intégrité, 65

- groupe de serveurs ILB, 61

- règles ILB, 68

- routeur VRRP , 38

- VNIC VRRP, 37

D

- démons

- in.ripngd démon, 22, 23

- désactivation

- routeur VRRP, 41

- dladm, commande

- create-vnic, 37

E

- /etc/inet/ndpd.conf fichier, 24

- création, 24

- équilibreur de charge intégré Voir ILB

- Ethernet via InfiniBand

- VRRP et, 33

G

groupe de serveurs ILB

affichage, 64

ajouter, 61

création, 61

suppression, 64

groupes de serveurs ILB

définition, 61

H

haute disponibilité

topologie DSR, 75

topologie Half-NAT, 78

I

ILB

activation, 59

affichage

statistiques, 72

table de connexions NAT, 72

table de mise en correspondance de la persistance
de session, 73

algorithmes, 68

aperçu, 13

autorisation d'utilisateur, 58

cas d'utilisation pour configurer un ILB, 70

composants, 49

contrôle de l'intégrité, 64

désactivation, 60

détails du test, 66

exemple de création d'un groupe de serveurs ILB et
de l'ajout de serveurs d'arrière-plan, 62exemple de la désactivation et de la réactivation de
serveur d'arrière-plan dans un groupe de serveurs
ILB, 63exemple de la suppression d'un serveur d'arrière-
plan d'un groupe de serveurs ILB, 64

exemple la création d'une règle full-NAT, 69

exportation

configuration, 74

fonctionnalités, 13

gestion, 60

groupes de serveurs, 61

haute disponibilité, 75, 78

importation

configuration, 74

installation, 57

ligne de commande, 58

mode DSR, 50

mode NAT, 50

modes de fonctionnement, 50

processus, 54

règles, 67

serveurs d'arrière-plan, 63

sous-commandes d'affichage, 59

sous-commandes de configuration, 58

statistiques

affichage, 71

in.ripngd démon, 22, 23

in.routed démon

mode économie d'espace, 11

in.routed, démon

Description, 10

installation

ILB, 57

VRRP, 36

ipadm command

create-addr, 40

IPv6

in.ripngd démon, 22

publication de routeur, 22

M

messages

publication de routeur, 22

messages Gratuitous ARP et et NDP, 45

mode direct server return *Voir* mode DSR

mode DSR

avantages, 50

description, 50

inconvénients, 50

mode économie d'espace

in.routed démon option, 11

mode NAT

avantages, 52

description, 51

inconvénients, 52

mode network address translator Voir mode NAT

modification

 routeur VRRP, 42

N

ndpd.conf fichier

 création, sur un routeur IPv6, 24

nouvelles fonctionnalités

 routeadm commande, 23

O

OSPF, 11

P

préfixe de site, IPv6

 publication, sur le routeur, 24

préfixes

 publication de routeur, 22

Protocole d'information de routage (RIP)

 Description, 10

Protocole de détection de routeur ICMP (RDISC), 11

protocole de routage

 VRRP, 12

protocole de routage, suite, 11

protocoles de routage

 BGP, 11

 démons de routage associés, 10

 description, 10

 OSPF, 11

 RDISC

 description, 11

 RIPng, 11

Protocoles de routage

 RIP

 Description, 10

publication du routeur

 IPv6, 22

Q

-q option

in.routed daemon, 11

R

RDISC

 description, 11

règles ILB

 création, 68

 liste, 68, 70

 suppression, 70

RIPng, 11

routeadm commande

 configuration de routeur IPv6, 23

routers

 aperçu, 9

routeur

 exemple de la configuration d'un routeur par défaut

 pour un réseau, 19

routeur IPv4

 configuration, 17

routeur IPv6

 configuration, 22

routeur VRRP

 activation, 41

 affichage de la configuration, 42

 affichage des adresses IP associées, 44

 aperçu, 12

 cas d'emploi pour la configuration d'un routeur

 VRRP, 46

 configuration de l'adresse IP virtuelle, 40

 création, 38

 Exemple d'affichage des informations de configuration de routeur de couche 3 au sein d'un système, 44

 exemple de configuration de l'adresse IP virtuelle pour un routeur VRRP de couche 3, 41

 exemple de l'affichage d'une adresse IP associée, 44

 exemple de la configuration d'un routeur VRRP de couche 3, 39

 exemple de la configuration d'une adresse IP virtuelle pour un routeur, 41

 exemple de la création d'un routeur VRRP, 39

 exemples de l'affichage des informations de configuration, 42

 modification, 42

 suppression, 45

routeurs VRRP et équilibreurs de charge
utilité, 15

routeurs
BGP, 11
configuration, 10
IPv6, 23
définition, 10
OSPF, 11
protocoles de routage
description, 10
quagga protocole de routage, suite, 11
RIPng, 11
VRRP, 12

S

-S option
in . routed démon, 11
serveur arrière-plan
réactivation, 62
serveur-à-client, 54
serveurs d'arrière-plan
suppression, 63
suppression
routeur VRRP, 45

T

Tables de routage
in . routed, création de démon de, 10
tables de routage
mode économie d'espace, 11
topologie
DSR, 50
Full-NAT, 54
Half-NAT, 53
topologie DSR
configuration, 75
topologie Half-NAT
configuration, 78

V

VNIC VRRP, 37

VRRP, 27
autorisation, 36
comparaison de couche 2 VRRP avec couche 3 VRRP, 31
configuration, 36
création d'une VNIC, 37
désactivation du routeur, 41
description, 12
installation, 36
inter-opérations
autres fonctionnalités du réseau, 32
limitations, 32
planification, 35
présentation, 27
Prise en charge d'Ethernet via InfiniBand, 33
prise en charge de zone en mode IP exclusif, 32
routeur de secours, 28
routeur maître, 28
VRRP couche 3
contrôle des messages Gratuitous ARP et NDP, 45
VRRP de couche 2
limitations, 32
vrrpadm command
show-router, 42
vrrpadm, commande
create-router, 36