

Data Deduplication in Windows Server 2012 R2 in Depth

Mike Ressler

Veeam Product Strategy Specialist, MVP, Microsoft
Certified IT Professional, MCSA, MCTS, MCP

Modern Data Protection
Built for Virtualization | **#1 VM Backup**

Introduction

Data deduplication has been around for years and comes in many forms. Most people think about expensive storage arrays when they talk about deduplication but that is certainly not the only type of data deduplication technology.

This white paper examines the data deduplication technology embedded in Windows Server 2012 R2 and the benefits it can bring to your environment. The discussion includes how it works and the possible scenarios for using Windows Server 2012 R2 data deduplication

Forms of data deduplication

Data deduplication is the technique used to eliminate duplicate or redundant information. A few methods can be used to achieve this.

Inline versus post-process

The first method concerns the timing of the data deduplication process. Inline deduplicates the data at the exact moment it enters the storage device. In other words, when a block of data is written to the storage device, it is checked immediately to see if this block already exists, and if it exists, it will not be stored but will reference the existing block. Post-process works the other way around. The block of data is stored first and at a later time the block will be analyzed.

Both methods have advantages and disadvantages. Inline ensures that you will need less storage. After all, when the data is deduplicated BEFORE it is actually written away, you will need less storage from the outset. The post-process method requires that you write the data in full first, before you can actually deduplicate it. This means that you need more storage at first, before you start saving data.

However, the inline method does have a disadvantage. Because it calculates at the moment of writing, you will need the resources to do this. And it can mean that the writing of your data is slower on the device, therefore reducing the backup throughput. The post-process method doesn't have this problem because it works in the background at certain intervals.

Source versus target deduplication

The second method concerns the location for data storage. With source deduplication, the data is deduplicated on the data source. In most cases, this refers to a location within the file system where the data lives. Target deduplication is the process of removing duplicates of data in the secondary store. This method is generally used with backup systems (such as a data repository or VTL) or replication targets.

File-, block- or chunk-based deduplication

File-based deduplication examines the data on a file level. For example, deduplication will occur when two Word documents are exactly the same. But when one of the two Word documents is changed slightly, the deduplication is undone. This method doesn't provide you with the best results.

Another method that is used a lot in storage arrays is block-based deduplication. Based on a fixed block size (e.g., 32 KB), the system looks at similarities in the blocks on the storage system and removes duplicate blocks with a reference. In this case, the smaller the block size, the better the deduplication because you will have more opportunity to find the same blocks. On the other hand, the smaller the block size, the more processing it takes to duplicate.

Chunk-based deduplication or chunking is comparable to block-based deduplication; the main difference is that instead of a fixed block size, it uses variable block sizes depending on what the algorithm determines as the best size at that moment.

Windows Server 2012 R2 data deduplication architecture

Before we dive into the architecture of data deduplication, it is important to know what kind of methods Microsoft uses for its implementation of data deduplication.

Microsoft uses the post-process, source and chunk-based methods. The downside is that you need the storage *before* the data is analyzed but you will need fewer resources to achieve your deduplication results. Because it is source-based, deduplication takes place on the storage where the actual data resides and, finally, the chunking method ensures that you achieve the best results.

The deduplication works on top of the NTFS file system and is a role that you can enable. Once enabled, it needs to be *activated* per **volume**.

There are a few requirements:

- The volume cannot be a boot or system volume
- MBR¹ or GPT² are supported but it needs the NTFS file system
- Deduplication cannot be enabled on CSVs³ (except for VDIs⁴)
- Remotely mapped drives are not supported

1. MBR: Master Boot Record

2. GPT: GUID Partition table

3. CSV: Cluster Shared Volume

4. VDI: Virtual Desktop Infrastructure

The deduplication engine uses a filter driver that lies on top of the NTFS file system. The filter driver is responsible for the coordination between the file entry, the regular storage and the chunk storage. Deduplication also adds a service to the operating system called *Data Deduplication Service (ddpsvc)*. This service is responsible for running the dedupe jobs at regular intervals (discussed later). There is also a service added called *Data Deduplication Volume Shadow Copy Service (ddpvssvc)*, which is the VSS writer for this functionality.

The *Data Deduplication Service* initiates three types of jobs (discussed later):

- Garbage collection
- Optimization
- Scrubbing

The variable chunks (for chunk-based deduplication) will be between 32 and 128 KB (averaging around 80 KB but this depends on your data.)

You should have the same requirements for your server hardware as Windows Server 2012 (R2); namely, it will run on a single processor system with 4 GB of RAM and one SATA drive. But when sizing your server you need to take into account a few considerations.

If you plan to support deduplication on multiple volumes on the same server, make sure that the system can process the data.

A server needs 1 CPU core and 350 MB of free memory to run a deduplication job on a single volume. That single job can process about 100 GB per hour and about 2 TB per day. When you use additional CPU cores and more available memory, that process can scale to enable parallel processing of multiple volumes.

The following example is taken from [TechNet](#) and shows how you can size your environment:

If you have a server with 16 CPU core processors and 16 GB of memory, deduplication uses 25% of the system memory in the default background processing mode, so in this case, that would be 4 GB. Dividing by 350 MB (free memory required per volume), you can calculate that the server could process about 11 volumes at a time. If you add 8 GB of memory, the system could process 17 volumes at a time. If you set an optimization job to run in throughput mode, the system would use up to 50% of the system's memory for the optimization job.

Data deduplication supports up to 90 volumes at a time; however, deduplication can simultaneously process one volume per physical CPU core processor plus one. Hyper-threading has no impact because only physical core processors can be used to process a volume. A system with 16 CPU core processors and 90 volumes will process 17 volumes at a time until all 90 volumes are done, if there is sufficient memory.

Virtual server instances should follow the same guidelines as physical hardware regarding server resources.

The last thing to explain is what happens when a file rehydrates. This refers to the fact that the file is taken out of deduplication and is kind of *rebuilt*. This is the process that occurs when you unoptimize your files (see “The jobs” section) and the file becomes a normal, non-deduped file. While you can do this on an entire volume by disabling the deduplication on that volume, you can also do this on a single file using the PowerShell cmdlets `Expand-DedupFile` for a single file.

Installation

When you want to enable deduplication on a Windows Server 2012 R2, you need to enable it through the GUI by adding the role under File and Storage Services (Figure 1).

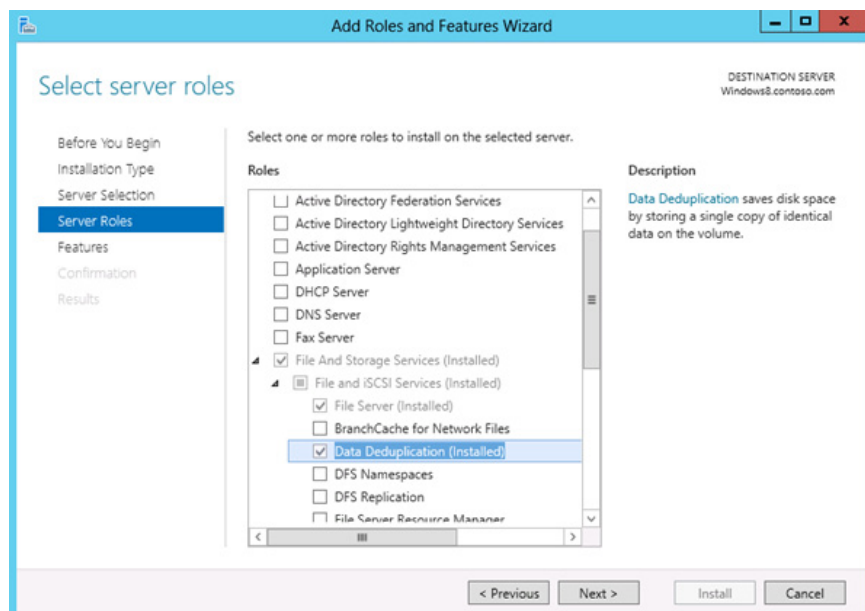


Figure 1: Adding data deduplication through the GUI

This can also be done through PowerShell, as shown below:

Import-Module ServerManager

Add-WindowsFeature -name FS-Data-Deduplication

Import-Module Deduplication

So far, the only thing you have done is to enable the role in Windows Server. Next you need to enable deduplication on a per-volume basis, either by using the GUI again or using PowerShell.

Through Server Manager, go to **File and Storage Services > Volumes** and choose the volume. Right-click and choose **Configure Data Deduplication...** (Figure 2).

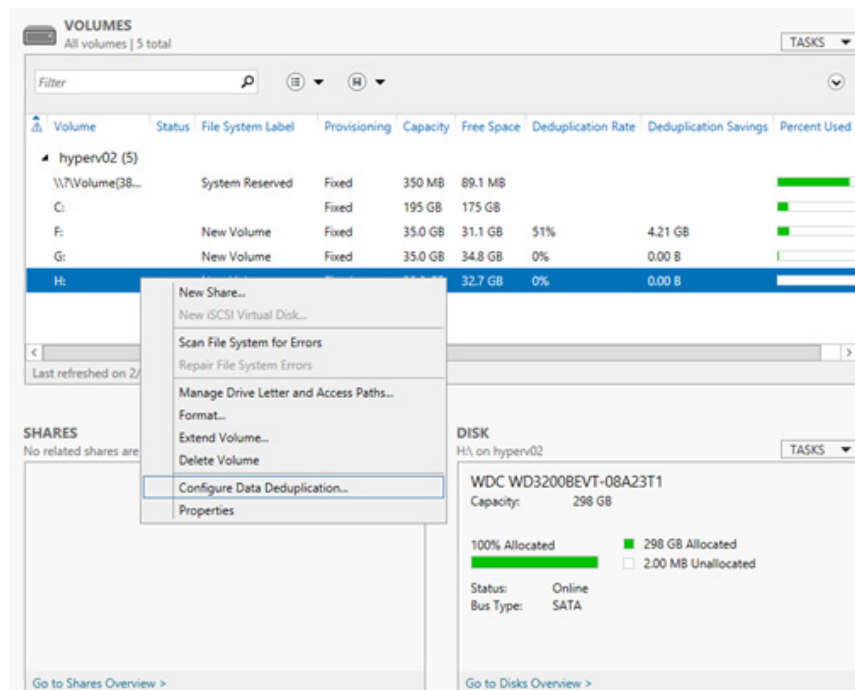


Figure 2: Enabling a volume through the GUI

This launches a new window that allows you to choose a few options for your volume (Figure 3).

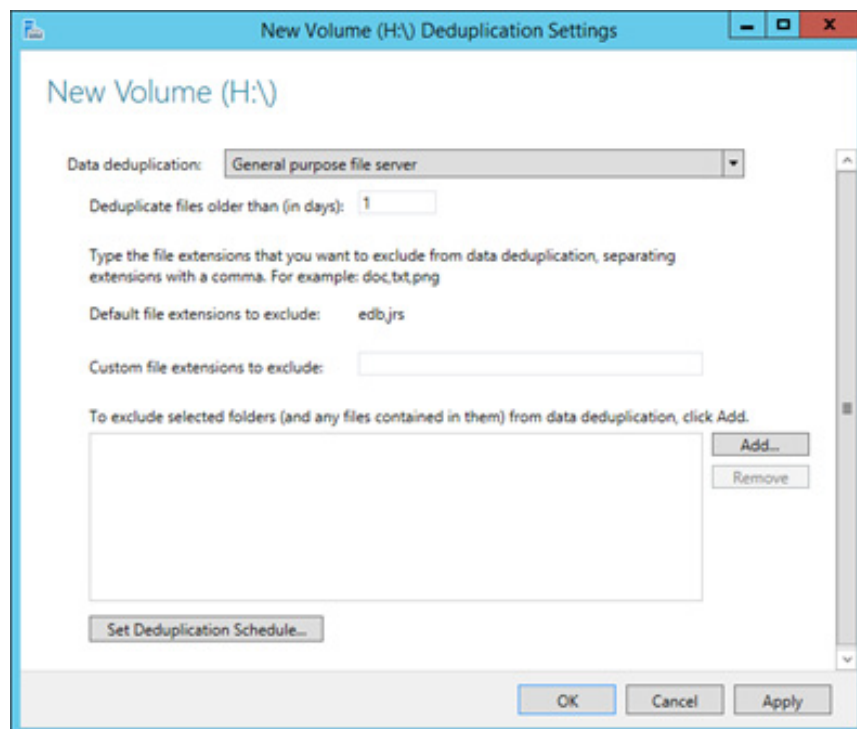


Figure 3: Data deduplication options for a volume

You can also do this using PowerShell. The following commands will enable the volume and set the deduplication files older than (in days) policy. Other settings can be configured through PowerShell as well:

Enable-DedupVolume H:

Set-DedupVolume H: -MinimumFileAgeDays 1

Now you have enabled data deduplication on your Windows Server and enabled it on a volume. From now on the deduplication will work.

The jobs

Data deduplication works through three jobs. This section examines what these three jobs do and how you can manually trigger them or optimize them through PowerShell.

Optimization

This job is responsible for deduplicating data and compressing the file chunks on a volume based on the policy settings per volume.

The job will run itself hourly but you can change a few settings or force it to run itself by using PowerShell:

Start-DedupJob -Volume H: -Type Optimization

By using the following command you can query the progress of the job:

Get-DedupJob

You can also look at the key statistics (free space, space saved, optimized files and more) by using the following command:

Get-DedupStatus | Format-List

Or you can view it through Server Manager and look at your volume.

More information about this job and about metadata that is collected can be found on [TechNet](#)

Garbage collection

Garbage jobs are responsible for processing deleted or modified data on the volume. With this job, data chunks that aren't referenced anymore are cleaned up. This means that when you delete data on the volume, you won't see savings immediately—you will need to wait until this job runs before it is actually cleaned up. Of course you can run this job manually to save disk space immediately, but keep in mind that this is a processing-intensive operation, so it might not be a smart move to run it during the production hours.

You can use the following PowerShell command to remove unreferenced chunks and compact containers that have more than 5% of unreferenced data:

Start-DedupJob H: -Type GarbageCollection

Adding the **-full** parameter to the PowerShell command ensures that all the containers are compacted to the maximum:

Start-DedupJob H: -Type GarbageCollection -full

Scrubbing

Data deduplication includes a weekly scrubbing job to make sure that your data remains integer. There are many reasons why data can encounter corruption but the scrubbing job ensures that any corruption is recorded in a log file:

Event Viewer\Applications and Services Logs\Microsoft\Windows\
Deduplication\Scrubbing

And when possible, it will make the repairs. There are a few repair options:

- **Hotspots:** The system keeps backup copies of popular chunks when they are referenced over 100 times (default). If the working copy is corrupted, deduplication will use that backup.
- **Storage Spaces:** When using deduplication on storage spaces in a mirrored configuration, deduplication can use the mirror image to fix the corruption.
- **Processing:** If a file is processed with a chunk that is corrupted, the corrupted chunk is eliminated and the new incoming chunk is used to fix the corruption.

The scrubbing job is launched on a weekly basis, but you can trigger this on demand by using PowerShell. This is again a processing-intensive job, so it is not advised to run this during production hours unless necessary:

Start-DedupJob E: -Type Scrubbing

If you want to do **deep scrubbing**, you can use the **-full** parameter. This ensures that repairs will be attempted on the entire set of deduplicated data, not just on the corruptions that are logged:

Start-DedupJob E: -Type Scrubbing -full

How it works

This section examines in more detail what happens when files get processed. The diagram below presents a use case for two different files (or file types) (Figure 4).

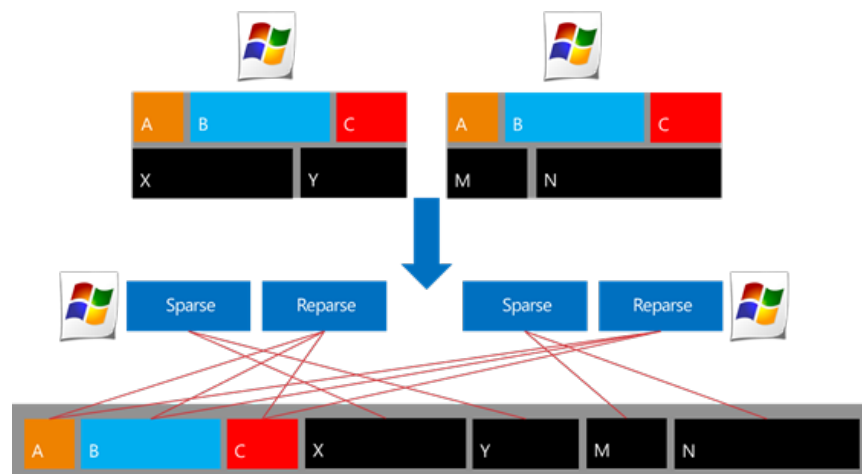


Figure 4: Model showing how files are processed

The two files are processed and duplicate chunks are found (A, B and C). As you can see, the length of these chunks is different. Each file also contains chunks that aren't duplicates (X, Y, M and N). The chunks are now stored in the System Volume Information folder on that volume. Note that in addition to the duplicate chunks that are stored there, the non-duplicate chunks will also be stored there. After that operation, the two files will be shown as 4 KB on their location. This is the size of the metadata (name, attributes, sparse and reparse data).

After the optimization job, the volume will contain:

- **Unoptimized files:** Files that do not meet the policy (explained later) or files that are excluded from the deduplication jobs
- **Optimized files:** Files that are stored as reparse points that contain pointers to a map of the chunks in the chunk store
- **Chunk store:** The location for the optimized file data

Policies / settings

There are quite a few policies that you can change to optimize deduplication in your environment and to achieve the best possible results.

Minimum file age: This setting allows you to specify the minimum file age in days based on the last modification time. However, please note that if Last Access Time is enabled on the server, the deduplication solution will use that instead.

Background mode: By default the optimization job runs every hour in background mode. This means that the system uses up to a maximum of 25% of the system memory. When you run an optimization job manually it will be able to use up to 50% of the available memory. On the deduplication schedule you can create two schedules to run this in throughput mode (50%) also.

Excludes: You can exclude file types or entire folders from being deduplicated. By default, in Windows Server 2012 R2, the file types EDB and JRS are excluded. In Windows Server 2012 nothing is excluded.

Garbage collection: The garbage collection job is scheduled at 1:45 a.m. on Saturday. If this conflicts with other maintenance jobs you should change this schedule as it is process-intensive.

Scrubbing: The scrubbing job runs at 2:45 a.m. on Saturday. This should be changed also if you have maintenance jobs running.

Size less than 32 KB: All files smaller than 32 KB will be ignored by default.

Most of these settings can be adapted through the GUI and all of them can be adapted by using the PowerShell command: Set-DedupVolume

PowerShell

This white paper has already introduced quite a bit of PowerShell. There is a simple command to find all the PowerShell cmdlets supported by deduplication (Figure 5):

Get-Command -Module Deduplication

```

PS C:\Users\administrator.MD> Get-Command -Module deduplication
CommandType      Name                                     ModuleName
-----
Function         Disable-DedupVolume                    Deduplication
Function         Enable-DedupVolume                    Deduplication
Function         Expand-DedupFile                      Deduplication
Function         Get-DedupJob                          Deduplication
Function         Get-DedupMetadata                    Deduplication
Function         Get-DedupSchedule                    Deduplication
Function         Get-DedupStatus                      Deduplication
Function         Get-DedupVolume                      Deduplication
Function         Measure-DedupFileMetadata            Deduplication
Function         New-DedupSchedule                    Deduplication
Function         Remove-DedupSchedule                 Deduplication
Function         Set-DedupSchedule                    Deduplication
Function         Set-DedupVolume                      Deduplication
Function         Start-DedupJob                      Deduplication
Function         Stop-DedupJob                        Deduplication
Function         Update-DedupStatus                    Deduplication
  
```

Figure 5: Data deduplication PowerShell cmdlets

VDI

Starting from Windows Server 2012 R2, data deduplication now supports virtual desktop infrastructure (VDI) as well. This means that you can dedupe your existing VDI machines and save tons of disk space.

By optimizing your CSV volumes for your VDI workloads you can have storage savings of up to 95%, which is huge! Most of the time, storage in a VDI implementation is the most expensive part of the implementation and with data deduplication, huge cost-savings can be achieved.

In order to achieve these savings, your VDIs need to run on a CSV with a scale-out file server (Figure 6).

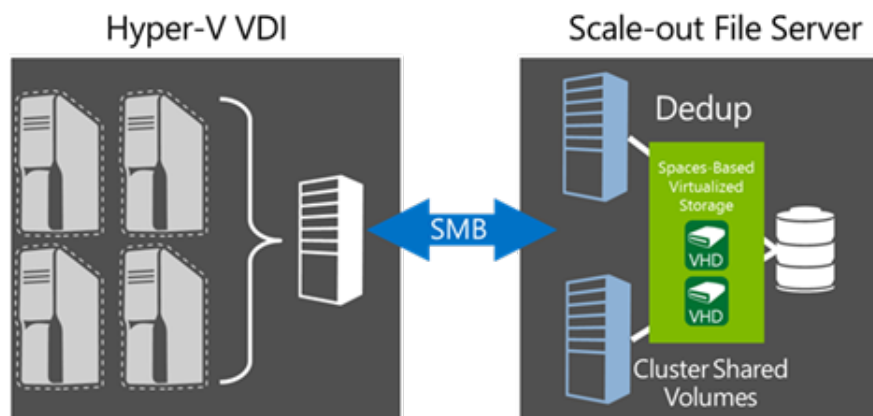


Figure 6: Hyper-V VDI requirements for deduplication; source: <http://blogs.technet.com/b/filecab/archive/2013/07/31/extending-data-deduplication-to-new-workloads-in-windows-server-2012-r2.aspx>

Besides saving space, you will also improve your I/O. Because data deduplication uses a caching mechanism that is optimized for VDI environments, your VDIs will run faster (the next paragraph explains how to enable deduplication on a CSV with VDIs). You will notice that your VDIs run faster during the daily boot storm. A lot of the same files are loaded at the same time (after all, a booting client system uses the same files) and the chunks for those files will be cached and therefore start faster.

To enable deduplication on a CSV with VDIs you need to tell the system to do so. You can do this through the GUI again by choosing the Virtual Desktop Infrastructure (VDI) server setting in your settings (Figure 7):

New Volume (H:\)

Data deduplication: Virtual Desktop Infrastructure (VDI) server

Figure 7: VDI setting for data deduplication

The PowerShell command looks like this:

Enable-DedupVolume C:\ClusterStorage\Volume1 -Usagetype Hyper-V

Benefits

How many benefits you will have depends entirely on your infrastructure, the type of files and the data in those files. Microsoft has released a guidance of savings that you should be able to get. Of course, this still depends on your data, but it should provide you with a reference for your environment. Typical deduplication savings for various content types are shown in Table 1:

Scenario	Content	Typical Space Savings
User documents	Documents, photos, music, videos	30-50%
Deployment shares	Software binaries, cab files, symbols files	70-80%
Virtualization libraries	Virtual hard disk files	80-95%
General file share	All of the above	50-60%

Table 1: Deduplication savings: Source: <http://technet.microsoft.com/en-us/library/hh831700.aspx>

Don't forget that the savings on the virtualization libraries are on non-running virtual machines (VMs) such as backups, stored VHDs, WIM files and so on. So don't use this on your running VMs.

Reliability and performance

One of the biggest fears with data deduplication is reliability. Imagine that you have a certain chunk that is used by more than a hundred files on your file server and that specific chunk gets corrupted—then you end up losing more than a hundred files. The data deduplication implementation of Microsoft has developed many control mechanisms to prevent corruption. We have already discussed one of these mechanisms—the scrubbing process. Other mechanisms are full checksum validation, redundancy for metadata and so-called hotspots for data chunks that are frequently referenced (e.g., when a hotspot is detected, a duplicate of the data chunk is created. By default, a chunk becomes a hotspot when it is referenced 100 times, but that can be changed through PowerShell).

As you can see, Microsoft has implemented quite a few techniques to make the solution reliable, but as always, you need to make sure that your data is being protected.

Before using your data protection solution, make sure that you know that data deduplication as it is implemented in Windows Server 2012 (R2) is supported by your vendor. Veeam Backup & Replication™ v7 R2 is one of the solutions that can protect a data deduplication-enabled volume.

Another big fear administrators have is the loss of performance. While this certainly can happen (e.g., when a file is accessed that is deduplicated it needs to be rehydrated), the performance loss is minimal with this technology.

The performance hit on writing data will be zero because it is a postponed process. However, there is also a performance hit on reading the data. You should expect around 3% performance loss in your environment.

However, since the solution works with a caching mechanism, it will store *frequently used chunks* in its memory and therefore it isn't unusual to see performance increases on file servers.

Interoperability

Data deduplication can work together with other Microsoft-related technologies, and this is another reason to investigate the solution.

BranchCache

If you are using BranchCache (a feature of Microsoft Windows) in your environment, it can benefit from data deduplication. When a server with BranchCache talks to a remote file server over the WAN, all the deduplication files are already indexed and hashed. This results in faster computation of branch office requests and greater probability of sending less data over the wire.

Failover clusters

Failover clusters are fully supported, and deduplicated volumes can fail over without any issue. Keep in mind that each node must be running the data deduplication feature.

DFS Replication

Many people believe that DFS and data deduplication can't work together. However, when you optimize or unoptimize a file, this will NOT trigger a replication because the file itself doesn't change. Since DFS uses Remote Differential Compression (RDC) and not the chunks in the data store, it will work together with data deduplication perfectly.

Tips before starting

Before you actually begin with data deduplication, you need to plan very carefully. Certainly in already existing environments, enabling deduplication on a volume without careful consideration can have bad consequences.

There are a few questions you need to ask yourself:

Can the data be deduplicated?

Some configurations are not supported, such as Single Instance Storage (SIS) or FRSM hard quotas. Other file types are just very difficult to deduplicate and don't bring in any advantages. Make sure you check (upfront!) with the [DDPEval](#) tool to estimate your potential savings. Other configurations that are not supported are constantly open or changing files (SQL or Exchange databases, for example). Additional structured data types that may have their own compression or single write instance technique may also not be good candidates for running data deduplication from Windows Server, as there would be unnecessary memory and I/O consumption for minimal benefit. While there is an exception for VDI environments, this is the only exception. Thus, server VMs are also not supported (running).

Is the volume empty?

It can be very *challenging* to enable deduplication on an existing volume. However, don't forget that you need to be able to finish a first *run*. When the volume contains a lot of data that is ready for deduplication, you may not be able to finish a run within the foreseen timeframe, or if your server is very active, it may respond slower. For these reasons, don't do this during business hours and do monitor to see that the first pass will finish. Use the throughput method to have more resources available for the process and schedule runs during the weekend or after hours.

Resources

Make sure you have enough resources on the server to support all the volumes that you want to deduplicate.

Plan your settings

Don't start by using the default settings when you know they aren't suited for your environment. Change the schedule of the garbage collection and data scrubbing when they interfere with other jobs (such as weekly backup jobs). Also, when changing the minimum file age (by default, three days) to fewer days, make sure that the job will be able to complete. It's better to start with a larger minimum file age and change it afterward to a lower number. And last but not least, don't forget exclusions by file type or folder.

Data protection

Ask your data protection vendor upfront whether they support data deduplication and how their solution works with it. Will their solution rehydrate the deduplicated files completely during the backup? Will it store the files in an optimized format on your backup target?

Size

While there is not an official limit for the amount of data that deduplication can handle, there are some issues. Volumes with files over approximately 10 TB have issues being deduped and probably will never succeed in their timeframe, causing them go into an infinite loop. If you have files of that size, you may run into this problem and then you will lose resources on your server for no reason. When necessary, exclude those type of files.

Using a Windows Server 2012 R2 deduplication-enabled volume as a Veeam repository

Veeam Backup & Replication provides both compression and deduplication features out of the box. Certainly in data protection, storage demands can become very high and grow much higher than the actual demand on the production storage.

Your mileage will vary

A number of factors go into data deduplication performance and scalability. For Veeam repositories, factors such as how much throughput is available to the repository makes a difference in how long it will take for the data deduplication tasks to complete. For example, the highest performing repository configurations are physical proxies that are also repositories running on locally attached JBOD storage. This environment would process through the Windows Server deduplication tasks very efficiently compared to a Microsoft Hyper-V or VMware vSphere VM with a large virtual disk format (VMDK or VHD or VHDX) that is on a storage system shared with other VMs.

Veeam Backup & Replication uses what is known as a “per job” deduplication, meaning that this deduplication is achieved per *backup file*. The results from this are already pretty impressive, saving you data storage costs by providing effectively global deduplication of your Veeam backup jobs with built-in features from Windows Server.

But what if we can take this one step further?

If we would create a backup repository on a Windows Server 2012 R2 volume and enable deduplication on it, we could save even more.

Using Veeam Backup & Replication v7 R2 and Windows Server 2012 R2, and combining both solutions (Veeam deduplication and compression and Windows Server 2012 R2 deduplication), you can achieve great savings on your storage.

You can find much more information on the blog written by [Chris Henley](http://www.veeam.com/blog/how-to-get-unbelievable-deduplication-results-with-windows-server-2012-and-veeam-backup-replication.html): <http://www.veeam.com/blog/how-to-get-unbelievable-deduplication-results-with-windows-server-2012-and-veeam-backup-replication.html>

Conclusion

If you haven't looked into data deduplication in Windows Server 2012 (R2), you should. The savings are huge, the performance is very good and most of all, it works right out of the box with your Windows Server license.

In addition to using deduplication for file servers, operating system deployment files, static content providers and more, it is also something to look for when you are using Veeam Backup & Replication and you are using a Windows Server 2012 (R2) server as your repository.

About the Author



Mike Ressler is a Product Strategy Specialist for Veeam. Mike is focused on technologies around Hyper-V and System Center. With years of experience in the field, he presents on many occasions at large events such as MMS, TechEd and TechDays. Mike has been awarded the MVP for System Center Cloud and Datacenter Management since 2010. His major hobby is discussing and developing solid disaster recovery scenarios. Additionally, he has enterprise-class experience in private cloud architecture and deployment, with marked focus on protection from the bottom to the top. He holds certifications in many Microsoft Technologies such as MCITP.

Follow Mike on [@MikeRessler](#) or [@Veeam](#) and on [Google+](#).

About Veeam Software

Veeam® is Protection for the Modern Data Center™ - providing powerful, easy-to-use and affordable solutions that are Built for Virtualization™ and the Cloud. [Veeam Backup & Replication™](#) delivers [VMware vSphere backup](#), [Hyper-V backup](#), recovery and replication. This #1 VM Backup™ solution helps organizations meet RPOs and RTOs, save time, eliminate risks and dramatically reduce capital and operational costs. [Veeam Backup Management Suite™](#) provides all the benefits and features of Veeam Backup & Replication along with advanced monitoring, reporting and capacity planning for the backup infrastructure. [Veeam Management Pack™](#) (MP) extends enterprise monitoring to vSphere through Microsoft System Center and also offers monitoring and reporting for the Veeam Backup & Replication infrastructure. The [Veeam Cloud Provider Program](#) (VCP) offers flexible monthly and perpetual licensing to meet the needs of hosting, managed service and cloud service providers. VCP currently has over 4,000 service provider participants worldwide. Monthly rental is available in more than 70 countries from more than 50 Veeam aggregators.

Founded in 2006, Veeam currently has 23,000 ProPartners and more than 91,500 customers worldwide. Veeam's global headquarters are located in Baar, Switzerland and the company has offices throughout the world. To learn more, visit <http://www.veeam.com>.



Protection for the
Modern Data Center



To learn more, visit <http://www.veeam.com/backup>