
Subject: Re: [ckrm-tech] [PATCH 0/2] resource control file system - aka containers on top of nsproxy!

Posted by [Paul Menage](#) on Thu, 08 Mar 2007 01:35:58 GMT

[View Forum Message](#) <> [Reply to Message](#)

On 3/7/07, Eric W. Biederman <ebiederm@xmission.com> wrote:

> Pretty much. For most of the other cases I think we are safe referring
> to them as resource controls or resource limits. I know that roughly covers
> what cpusets and beancounters and ckrm currently do.

Plus resource monitoring (which may often be a subset of resource control/limits).

>

> The real trick is that I believe these groupings are designed to be something
> you can setup on login and then not be able to switch out of.

That's going to be the case for most resource controllers - is that the case for namespaces? (e.g. can any task unshare say its mount namespace?)

Paul

Containers mailing list

Containers@lists.osdl.org

<https://lists.osdl.org/mailman/listinfo/containers>

Subject: Re: [ckrm-tech] [PATCH 0/2] resource control file system - aka containers on top of nsproxy!

Posted by [ebiederm](#) on Thu, 08 Mar 2007 02:25:54 GMT

[View Forum Message](#) <> [Reply to Message](#)

"Paul Menage" <menage@google.com> writes:

> On 3/7/07, Eric W. Biederman <ebiederm@xmission.com> wrote:

>> The real trick is that I believe these groupings are designed to be something
>> you can setup on login and then not be able to switch out of.

>

> That's going to be the case for most resource controllers - is that
> the case for namespaces? (e.g. can any task unshare say its mount
> namespace?)

With namespaces there are secondary issues with unsharing. Weird things like a simple unshare might allow you to replace /etc/shadow and thus mess up a suid root application.

Once people have worked through those secondary issues unsharing of

namespaces is likely allowable (for someone without CAP_SYS_ADMIN). Although if you pick the truly hierarchical namespaces the pid namespace unsharing will simply give you a parent of the current namespace.

For resource controls I expect unsharing is likely to be like the pid namespace. You might allow it but if you do you are forced to be a child and possible there will be hierarchy depth restrictions. Assuming you can implement hierarchical accounting without too much expense.

Eric

Containers mailing list
Containers@lists.osdl.org
<https://lists.osdl.org/mailman/listinfo/containers>

Subject: Re: [ckrm-tech] [PATCH 0/2] resource control file system - aka containers on top of nsproxy!

Posted by [Herbert Poetzl](#) on Fri, 09 Mar 2007 00:56:39 GMT

[View Forum Message](#) <> [Reply to Message](#)

On Wed, Mar 07, 2007 at 05:35:58PM -0800, Paul Menage wrote:

> On 3/7/07, Eric W. Biederman <ebiederm@xmission.com> wrote:

>> Pretty much. For most of the other cases I think we are safe

>> referring to them as resource controls or resource limits. I know

>> that roughly covers what cpusets and beancounters and ckrm currently

>> do.

>

> Plus resource monitoring (which may often be a subset of resource

> control/limits).

we (Linux-VServer) call that resource accounting, and it is the first step to resource limits ...

> > The real trick is that I believe these groupings are designed to be

> > something you can setup on login and then not be able to switch out

> > of.

>

> That's going to be the case for most resource controllers - is that

> the case for namespaces? (e.g. can any task unshare say its mount

> namespace?)

ATM, yes, and there is no real harm in doing so this would be a problem for resource containers,

unless they are strict hierarchical, i.e. only
allow to further restrict the existing resources
(which might cause some trouble if existing limits
have to be changed at some point)

best,
Herbert

> Paul

> _____

> Containers mailing list

> Containers@lists.osdl.org

> <https://lists.osdl.org/mailman/listinfo/containers>

Containers mailing list

Containers@lists.osdl.org

<https://lists.osdl.org/mailman/listinfo/containers>
