
Subject: Re: Re: [RFC] Virtualization steps
Posted by [Sam Vilain](#) on Tue, 28 Mar 2006 23:28:13 GMT
[View Forum Message](#) <> [Reply to Message](#)

On Wed, 2006-03-29 at 02:24 +0400, Kir Kolyshkin wrote:
> >Huh? You managed to measure it!? Or do you just mean "negligible" by
> >"1-2 per cent" ? :-)
> We run different tests to measure OpenVZ/Virtuozzo overhead, as we do
> care much for that stuff. I do not remember all the gory details at the
> moment, but I gave the correct numbers: "1-2 per cent or so".
>
> There are things such as networking (OpenVZ's venet device) overhead, a
> fair cpu scheduler overhead, something else.
>
> Why do you think it can not be measured? It either can be, or it is too
> low to be measured reliably (a fraction of a per cent or so).

Well, for instance the fair CPU scheduling overhead is so tiny it may as well not be there in the VServer patch. It's just a per-vserver TBF that feeds back into the priority (and hence timeslice length) of the process. ie, you get "CPU tokens" which deplete as processes in your vserver run and you either get a boost or a penalty depending on the level of the tokens in the bucket. This doesn't provide guarantees, but works well for many typical workloads. And once Herbert fixed the SMP cacheline problems in my code ;) it was pretty much full speed. That is, until you want it to sacrifice overall performance for enforcing limits.

How does your fair scheduler work? Do you just keep a runqueue for each vps?

To be honest, I've never needed to determine whether its overhead is 1% or 0.01%, it would just be a meaningless benchmark anyway :-). I know it's "good enough for me".

Sam.
