



ESCENA DIGITAL 2.0, REPOSITORIO DEL MUSEO DE LAS ARTES ESCÉNICAS (MAE)



Anna Valls y Roger Guasch



Anna Valls es licenciada en documentación por la *Universitat Oberta de Catalunya* y en filología catalana por la *Universitat de Barcelona*. Posgrado en gestión y planificación de bibliotecas de la *Universitat Pompeu Fabra* y en competencias directivas de la *Universitat Autònoma de Barcelona*. Desde 2006 dirige el MAE, *Centro de Documentación y Museo de las artes escénicas*. Durante 15 años trabajó en el ámbito de las bibliotecas universitarias.
<http://orcid.org/0000-0003-3891-7482>

vallspa@institutdelteatre.cat



Roger Guasch es ingeniero superior en informática en la *Facultad de Informática* de la *UPC*, especialización en ingeniería del software y sistemas de información, y en gestión y explotación de la información. Postgrado en desarrollo de proyectos de ingeniería y consultoría TIC y máster en dirección de sistemas de información (*UPC School of Professional & Executive Development*). Desde 2008 es responsable TIC del MAE, *Centro de Documentación y Museo de las Artes Escénicas*.
<http://orcid.org/0000-0003-0717-5973>

guaschar@institutdelteatre.cat

Centro de Documentación y Museo de las Artes Escénicas (MAE)
Institut del Teatre de Barcelona
Plaça Margarida Xirgu, s/n. 08004 Barcelona, España

Resumen

Bibliotecas, archivos y museos difunden sus fondos cada vez más a través de repositorios digitales online. Sin embargo, la evolución de estos repositorios está estancada desde hace años, tanto por su concepción funcional como por la tecnología asociada. Un grupo de universidades con larga tradición en este ámbito ha fundado el proyecto *Hydra*, que permite a cada institución crear su repositorio y dotarlo de las utilidades que desee. Se expone la experiencia del *Centro de Documentación y Museo de las Artes Escénicas (MAE)*, uno de los primeros socios europeos que han apostado por esta nueva generación de repositorios. Se explica qué y cómo se hizo, los problemas encontrados, los errores y los aciertos y finalmente las conclusiones obtenidas de todo el proceso.

Palabras clave

Repositorios digitales, Innovación, Proyecto *Hydra*, Difusión online, Artes escénicas, MAE, *Museo de las artes escénicas*, Barcelona.

Title: *Digital scene 2.0, repository of the Museum of the Performing Arts (MAE)*

Abstract

Libraries, archives and museums are increasingly committed to making their resources available through online digital repositories. However, progress in this area has been stalled for years, both in the functional design arena and in the relevant technology. Now, a group of universities with a long tradition in this area has created the *Hydra Project*: a framework that allows each institution to create their own repository and provides all needed utilities. This article describes the experience of the *Documentation Centre and Museum of the Performing Arts*, one of the first European partners to opt for this new generation of repositories, when faced with this challenge. It aims to explain what was done and how, the problems encountered, the mistakes and successes, and finally, the conclusions that can be drawn from the whole process.

Keywords

Digital repositories, Innovation, Hydra project, Online dissemination, Performing arts, MAE, *Museum of the Performing Arts*, Barcelona.

Valls, Anna; Guasch, Roger (2013). "Escena digital 2.0, repositorio del Museo de las Artes Escénicas (MAE)". *El profesional de la información*, 2013, mayo-junio, v. 22, n. 3, pp. 244-249.

<http://dx.doi.org/10.3145/epi.2013.may.08>

Artículo recibido el 14-02-2013
Aceptación definitiva: 06-05-2013

Introducción

El propósito de este artículo es presentar la experiencia del *Centro de Documentación y Museo de las Artes Escénicas* de Barcelona (MAE) al elaborar su repositorio digital usando el sistema *Hydra*. Se describe el entorno TIC en el inicio del proyecto, se exponen las alternativas estudiadas, se enumeran los aspectos técnicos, se apuntan costes asociados, aciertos y errores cometidos y se desglosan finalmente las conclusiones obtenidas.

Antecedentes

En 1996 el MAE se quedó sin un espacio de exhibición permanente. Diez años después, ante las malas perspectivas de conseguir otra sala de exposición, se decidió apostar por las TIC para potenciar la difusión de sus fondos. En 2007 se comenzó a fotografiar y escanear las piezas del museo y se implantó un primer repositorio digital usando el programa propietario *ContentDM 4.0*. La sencillez de su instalación y la velocidad de sus búsquedas contrastaba con su poco atractivo visual, la complejidad de su sistema de catalogación y la falta de algunas prestaciones elementales.

En 2009 se publicó la hemeroteca mediante *DSpace 1.5*. Aun teniendo una estética poco amigable resultaba una aplicación muy intuitiva, práctica y rápida. Lamentablemente su arquitectura interna era poco elegante: mezclaba la capa de presentación con la de negocio y ésta con la de datos. Para añadir más de los seis metadatos con los que trabajaba en un primer nivel, no solamente se debía reescribir parte de la configuración sino que se debía alterar el modelo de datos subyacente. Estas deficiencias lo hacían poco apto para mantener y evolucionar.

La problemática

En 2011 se evaluaron los procesos internos descubriendo tres aspectos a solucionar.

1) La ratio de catalogación era demasiado baja, porque:

- cada colección tenía un esquema de metadatos distinto: la falta de homogeneización dificultaba la posibilidad de tener validaciones automáticas, vocabulario controlado, procedimientos estándar, etc.;
- la aplicación de escritorio para catalogar en *ContentDM* no era intuitiva y padecía errores inexplicables.

2) El peso de los ficheros master es cada vez mayor: las fotografías *pesan* unos 25 MB, el escaneo de pósteres produce imágenes de hasta 500 MB y los vídeos en alta definición pueden alcanzar los 18 GB.

El público en general no puede acceder a estos archivos, pero buena parte de los usuarios son investigadores que necesitan consultarlos. Si se quería disponer de un sistema con funciones de preservación, catalogación y difusión, había que almacenar estos ficheros directamente en la aplicación. Pero el tiempo necesario para cargar esos ficheros era exasperante e incluso en ocasiones imposible (pocos repositorios aceptan subidas de elementos de más de 1 GB). Consecuentemente se acabó teniendo los ficheros master en soportes externos (discos duros, DVDs, etc.) sin integración con los registros web.

3) Los usuarios más exigentes echaban de menos ciertas funciones:

- crear, guardar, imprimir y exportar listados de resultados personalizados;
- que las búsquedas abarcaran más metadatos; sobre todo los más específicos de gestión interna y que en general no se compartirían mediante el protocolo OAI/PMH;
- que se pudieran reproducir nativamente diversos formatos multimedia.

Ante estas carencias se decidió desarrollar un sistema propio. Dadas las restricciones temporales y económicas se descartó implementar una aplicación desde

cero o usar un programa propietario. Quedaba la opción de probar algún repositorio de código abierto suficientemente dúctil para responder las necesidades planteadas, y se evaluaron distintas opciones.

Las alternativas

Se analizaron, excluyendo otras aplicaciones minoritarias, las tres grandes plataformas libres que acaparan casi el 70% del mercado: *EPrints*, *DSpace*, y *Fedora*.

Se descartó *EPrints* por la impresión de que está cayendo en desuso y sólo presenta actualizaciones menores de vez en cuando (*EPrints*, 2010). Comparando su utilización actual, con un 14% de cuota, con la de 2007 vimos que su porcentaje ha bajado 7 puntos. Además entre la versión 3.1 y la 3.3 -la última disponible- han pasado tres años.

El caso de *DSpace* (*The DSpace Foundation*, 2009) fue distinto. Se hizo un estudio por-

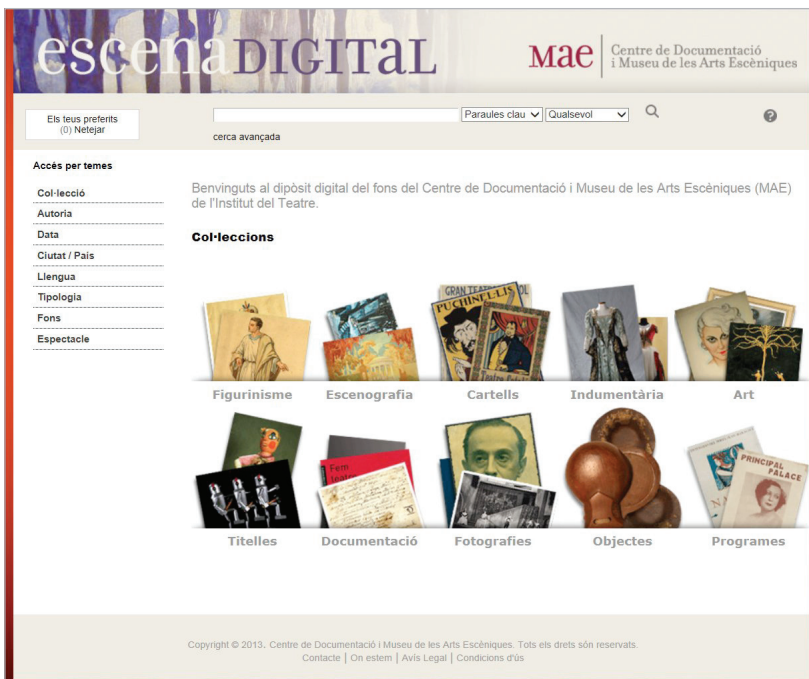


Figura 1. La nueva *Escena digital*
<http://colleccions.cdmae.cat>

menorizado porque era un producto que se conocía bien y acababa de sacar la versión 1.7. Sin embargo, se comprobó que sus problemas de arquitectura persistían. De haber optado por esta plataforma hubieran existido problemas con futuras evoluciones. Por ejemplo, si se quisiera:

- añadir nuevos metadatos en los que buscar (algo bastante probable después de las reuniones con los usuarios y demás interesados), habría que haber modificado el modelo de datos;
- dotar al sistema de funciones adicionales, habría que programarlas o aprovechar las de terceros. Y aun así, dichas implementaciones serían muy difíciles de mantener en futuras actualizaciones del propio *DSpace* al no poder desacoplarse claramente del núcleo de la aplicación.

“Hydra aunaba lo mejor de *ContentDM* y de *DSpace* pero sin las limitaciones de cada uno”

Fedora es un sistema de almacenaje de objetos digitales sin capa de presentación ni de catalogación (*Fedora Project*, 2007). Aunque la ayuda que suponía no era despreciable, obtener un sistema plenamente funcional implicaba dedicar excesivas horas de desarrollo.

La elección

Al hacer la comparación se vio que *DSpace* y *Fedora* dependen ambos de la organización *DuraSpace*, que junto con algunas universidades¹ habían creado un nuevo proyecto de código libre: *Hydra* (Dushay; Zumwalt, 2013; Green; Zumwalt, 2012; Green, 2013). Este proyecto nació para crear un producto que:

- a) contara con las funciones habituales;
- b) añadiera prestaciones que le dotaran de más valor añadido;
- c) no tuviera ninguna restricción que afectara a futuros desarrollos;
- d) no fuera necesario implementarlo desde cero.

Los responsables técnicos de las universidades llegaron a la conclusión de la necesidad de unir fuerzas y crear herramientas para que cada centro pudiera construir su repositorio de manera sencilla y autónoma. Era una idea que combinaba la flexibilidad con la ventaja de partir de un núcleo base. Además aportaba beneficios como:

- interfaz web atractiva y fácilmente personalizable;
- tecnología de vanguardia que garantizaba un buen rendimiento²;
- *plugins* y *widgets* que enriquecían el abanico de funciones;
- posibilidad de almacenamiento y reproducción de todo tipo de formatos multimedia;
- sistema basado en modelos flexibles que permiten construir modelos de datos ajustados a necesidades particulares;

- posibilidad de subida de ficheros asíncrona, que permite trabajar en un sistema parecido al de *Dropbox* con *uploads* masivos y pesos elevados.

Entre los inconvenientes:

- no se disponía de ningún profesional con conocimientos en *Ruby on Rails* (RoR);
- la comunidad detrás del proyecto era pequeña, de apenas diez miembros.

Aun sopesando dichas dificultades parecía que *Hydra* aunaba lo mejor de *ContentDM* y de *DSpace* pero sin las limitaciones de cada uno. Además las universidades fundadoras tenían un *know-how* amplísimo en repositorios digitales y pensamos que en el futuro nos enriquecerían con sus ideas.

El proyecto

Tecnología

Hydra se basa en cuatro grandes bloques³:

- *Fedora* como repositorio de objetos digitales;
- *Solr* como índice para las búsquedas;
- *Blacklight* (plugin RoR) para las funciones típicas del *Retrieve*;
- una gema *Ruby*⁴ llamada *HydraHead* que provee las funciones *create*, *update* y *delete*.

Modelo de datos

Hydra no exige que se use un único esquema. Puede haber tantos como se quiera, que deben estar definidos en la aplicación porque serán los modelos que usarán los objetos *Fedora*. Sí especifica que dichos objetos deben contener los siguientes *datastreams*⁵:

- *Rights metadata*: sencillo xml que contiene información sobre quién puede ver este contenido, quién editarlo, etc.;
- *DC metadata*: xml con los metadatos que describen el objeto;
- *RELS-EXT*: xml donde se expresan las relaciones entre este objeto y otros que pueden existir en el propio repositorio.

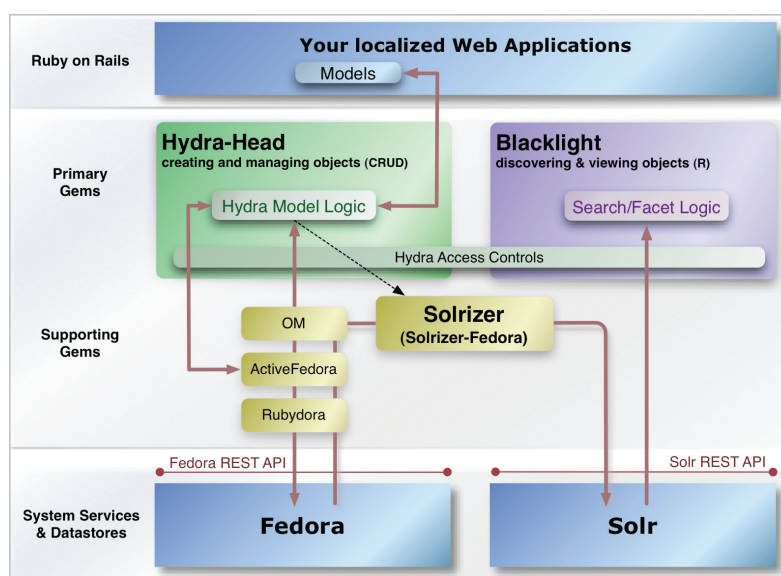


Figura 2. Herramientas e interacciones de Hydra (fuente: *DuraSpace*)

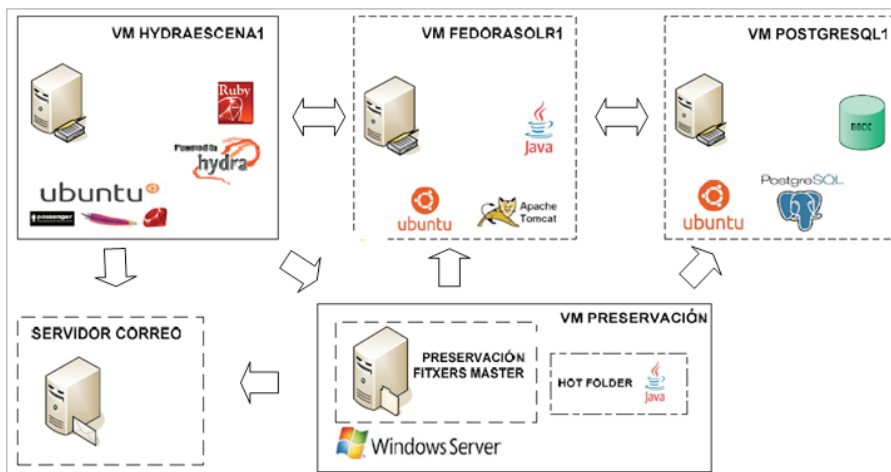


Figura 3. Servidores y componentes de *Escena digital* (fuente: MAE)

Mejoras realizadas

Aunque *Hydra* cubría la mayoría de nuestras necesidades, se implementaron algunas funciones que parecieron importantes.

1) Un modelo de seguridad más severo que incluye una gestión holista de usuarios, grupos, roles y permisos. La granularidad llega hasta el nivel de metadato, es decir, dependiendo del perfil del usuario, éste podrá ver unos datos u otros. Actualmente *Hydra* está trabajando en incorporar un sistema parecido en sus nuevas versiones.

2) Una gestión de vocabulario controlado en los metadatos. Dichos campos pueden ser identificados por el administrador en tiempo de ejecución. Una vez se determina que un atributo debe ser fiscalizado, se le aplicarán las reglas de negocio pertinentes para vigilar sus posibles valores. Así pues, a parte de las opciones más habituales como buscar, añadir, editar o eliminar, se permitirá hacer reemplazos masivos.

3) Una minuciosa auditoría del sistema. Se monitorea el trabajo que hace cada catalogador para saber qué y cuándo se ha hecho. Con ello se cubren tres objetivos:

- se minimizan los riesgos asociados a los cambios masivos ya que se informa inmediatamente a un responsable de las modificaciones acaecidas;
- se genera un histórico de versiones que puede ayudar a restablecer los datos correctos de un registro en caso de que una equivocación los hubiera cambiado o suprimido accidentalmente.
- obtener información estadística para los indicadores de productividad (KPIs, *key performance indicators*).

4) Creación de un sistema de ingesta de grandes archivos digitales a través de lo que se denomina “carpeta caliente”. Es una carpeta compartida de *Windows* que el sistema está monitorizando constantemente. Cuando detecta que se ha copiado un fichero, activa un disparador que copia este archivo en el depósito digital y lo asocia a un registro previamente catalogado. Este fichero master es accesible desde la plataforma web pero sólo para los catalogadores o investigadores. Una vez preservada dicha copia se genera automáticamente una versión en baja resolución que los usuarios pueden ver y descargar. La ganancia con este sistema es evi-

dente: las subidas son mucho más rápidas porque van por la red de fibra óptica interna y usan el protocolo *Samba*. Además están desacopladas del resto del sistema con lo que se pueden gestionar asincrónicamente sin sobrecargar la parte web.

Infraestructura

Uno de los aspectos clave para diseñar una aplicación con una arquitectura tan compleja es identificar la criticidad de sus módulos y cómo desacoplarlos para minimizar los puntos de fallo. Disponemos de tres servidores físicos, un NAS (*network-attached storage*)

conectado mediante *iSCSI* (*internet small computers system interface*) y el software *vSphere 4.1* de virtualización. Se decidió separar los componentes de la siguiente manera:

- aplicación *Hydra* y servidor web *Apache*;
- repositorio *Fedora* e índice *Solr*;
- base de datos *PostgreSQL*;
- sistema del *HotFolder* que preserva y produce las versiones de difusión de los objetos digitales.

“ La lista de requerimientos debe ser acotada o se corre el riesgo de que se pidan demasiadas funciones que acaben provocando retrasos en la entrega final ”

Esta configuración permite:

- Articular respuestas específicas en función de los problemas que puedan surgir. Por ejemplo, si hay un pico de peticiones concurrentes se incrementará la memoria y el ancho de banda disponible del servidor web; si nos proponemos hacer inserciones masivas de datos añadiremos CPU a la máquina virtual del *Fedora* y del *Solr*.
- Mejorar la seguridad del sistema, ya que la única parte accesible remotamente es la capa de presentación. Los demás elementos están detrás de un *firewall* y atendiendo peticiones únicamente de las máquinas con las que deben interactuar.
- En caso de pérdida total de una de las máquinas virtuales, al estar replicadas se puede levantar de nuevo en el mínimo tiempo sin afectación del resto de componentes.
- Soportar altas cargas de tráfico manteniendo unas óptimas prestaciones.

Rendimiento

Las pruebas de estrés a las que se sometió al sistema arrojaron resultados satisfactorios:

- Con 55.000 registros y 750 GB en imágenes, el tiempo medio de respuesta de una búsqueda en todo el catálogo demora unos 2,7 segundos.
- Con cerca de 5 millones de registros y unos 3 TB de imáge-

- nes, el tiempo medio de respuesta fue de 3,6 segundos.
- El tiempo medio de acceso a un registro es de 2,1 segundos cuando cuenta con una única imagen.
 - Para más imágenes, establecer una métrica exacta es más complejo ya que depende de su tamaño. No obstante las pruebas señalan que con 10 imágenes se tarda de media unos 6 segundos en cargar la página. En definitiva una función bastante aproximada sería calcular $2,1 \text{ segundos} + (0,4 \times \text{número de imágenes})$
 - Finalmente, el sistema soportó sin ninguna incidencia unos accesos concurrentes equivalentes a 50.000 usuarios y un total de 2.000.000 visitas en un sólo día. Estas cifras confirman que la escalabilidad del sistema es muy alta.

“Nunca hay que subestimar la curva de aprendizaje de una nueva tecnología, y siempre hay que hacer valoraciones temporales conservadoras”

Organización, tiempo y coste

Se construyó el nuevo repositorio en base a los requerimientos de los catalogadores. Para hacerlo había dos opciones: encargar el proyecto a una empresa o bien desarrollarlo con los ingenieros en plantilla. Ambas alternativas tienen ventajas e inconvenientes. La primera tiene un coste adicional muy alto y se pierde parte del *know-how* sobre la propia plataforma. En cambio libera recursos y es útil cuando el personal interno no tiene los conocimientos o habilidades necesarios para lograr un proyecto de este calaje. La segunda conlleva una paralización de otros proyectos pero los costes asociados son menores. Se optó por esta última.

Nuestra institución tiene un organigrama orientado a proyectos. Hay distintas áreas funcionales verticales y un departamento de informática transversal. Gracias a este modelo el responsable TIC se coordinaba con los jefes de cada área para compartir los recursos humanos. Así se maximizaba el tiempo efectivo de las personas y se mantuvo siempre una visión completa del estado del proyecto.

Este tipo de organización es la más pertinente en estos proyectos. Sólo con la participación activa de todos los usuarios internos se pueden alcanzar los objetivos propuestos. El éxito de cualquier proyecto TIC se mide por el uso de la herramienta que se está desarrollando. Si las personas no dan su opinión o no están convencidas de la utilidad del sistema, éste nunca podrá triunfar. Sin embargo la lista de requerimientos debe ser acotada o se corre el riesgo de que se pidan demasiadas funciones que acaben provocando retrasos en la entrega final.

Para evitar tales situaciones hubo que limitar las intervenciones, preparar cuestionarios cerrados, hacer reuniones con órdenes del día, redactar las actas, firmar los documentos funcionales, cuantificar temporal y económicamente cada requerimiento, etc. Actuar tan estrictamente provocó inicialmente ciertas tensiones: desacuerdos entre los documentalistas acerca de qué era prioritario, con los técnicos por el alcance del proyecto, con la dirección por su rigidez

en las fechas de entrega, etc. Posteriormente, cuando se nivelaron las expectativas, el trabajo en equipo resultó más fluido y acabó generando ideas muy buenas.

A la vista de la implicación requerida por parte de toda la plantilla es lógico preguntarse cuánto tiempo habrá que bloquear cualquier otra actividad. En nuestro caso tuvimos a un técnico con dedicación exclusiva y otro parcialmente durante los 13 meses que duró el proyecto:

- 2 meses en la toma de requisitos con los usuarios internos (catalogadores y jefes de sección) y la definición funcional;
- 1 mes para el estudio de mercado;
- 1 mes para el diseño técnico;
- 8 meses para el desarrollo;
- 2 meses para la migración de los datos (en paralelo al desarrollo);
- 1 mes para las pruebas.

En un proceso tan largo se cometen errores técnicos y funcionales. Entre los que se cometieron:

- Inicialmente todos los componentes estaban en un mismo servidor virtual. Como se ha explicado conllevaba muchos riesgos que se han atenuado al separarlos.
- Al principio se decidió almacenar los ficheros de preservación dentro de los objetos *Fedora*. El resultado fue que penalizaba enormemente el rendimiento. La solución fue guardar esos archivos en el NAS y añadir en un *datas-tream* un enlace hacia el fichero original.
- Para subir los ficheros se intentó aprovechar un componente típico de “examinar”, como se usa habitualmente en las plataformas web. Se demostró peligroso para la integridad de la aplicación ya que desbordaba la memoria del servidor y provocaba la finalización inesperada del sistema. La solución ideada se basa en la utilización de una carpeta *hot* como punto de entrada de los ficheros digitales.
- La coordinación en algunas fases del desarrollo no fue la adecuada. La consecuencia fue que después de hacer la migración de los datos se identificaran incongruencias que hubo que solventar con posterioridad.
- Aun habiendo firmado un documento de requerimientos a implementar se pidieron nuevas funciones una vez el producto estaba acabado.

Conclusiones

Cada organización es singular. Lo que puede funcionar en una quizá no resulte en otra. Después de años de experiencia se decidió que los repositorios más habituales no eran suficientes, pero no hubiéramos tomado la decisión de emprender el desarrollo de un repositorio *ad hoc* si no hubiésemos cumplido ciertas condiciones:

- Identificamos claramente nuestras debilidades y oportunidades.
- Nuestros catalogadores ya habían participado como *stakeholders* (interesados) en otros proyectos TIC y eso les dotaba de un *know-how* indispensable para abordar este reto.
- Nuestros ingenieros informáticos tenían gran experiencia en el sector de los repositorios digitales.

Resulta indispensable formalizar la dirección del proyecto. Se debe gestionar siguiendo una metodología más o menos estricta, como por ejemplo alguna simplificación del *Project Management Body of Knowledge (Pmbok)*⁶. Sin cierta rigidez el desvío temporal puede ser exagerado. Hay que tener presente que los usuarios internos, aunque participen en la definición del sistema, reparten su tiempo con el cumplimiento de sus tareas habituales, y eso puede desembocar en retrasos insalvables.

El sistema centraliza todo el proceso de gestión del fondo: ingesta, codificación, transformación, almacenamiento, catalogación y difusión online

El desconocimiento de la tecnología en la que se basa *Hydra* hizo que no calculáramos con precisión el tiempo que se tardaría en completar el proyecto. El resultado fue que, aun siendo generosos en las estimaciones, hubo que hacer infinidad de horas extras los últimos meses. Así pues, nunca hay que subestimar la curva de aprendizaje de una nueva tecnología. Y siempre hay que hacer valoraciones temporales conservadoras.

Con todo, la adopción del *framework Hydra* fue un acierto. No sólo hemos eludido las limitaciones que otros productos imponían sino que se ha construido un repositorio digital que satisface completamente los requerimientos de nuestros usuarios y que tiene una interfaz muy amigable. Además, después de las pruebas de rendimiento disponemos de un producto altamente escalable que ofrece unas prestaciones encomiables.

Con la nueva *Escena digital* podemos ofrecer un mejor y más amplio abanico de funciones online. Pero lo más destacable es que ahora disponemos de un sistema que centraliza todo el proceso de gestión de nuestro fondo: desde la ingesta, pasando por la codificación, transformación, almacenamiento y catalogación, hasta la difusión online. El efecto más evidente de trabajar en un entorno integral, ágil y sencillo, es que nuestros catalogadores maximizan su eficiencia.

En conclusión, gracias al nuevo repositorio el MAE dispone de más tiempo y recursos para invertir en servicios de valor

añadido que redunden en beneficio de sus usuarios.

Notas

1. *University of Virginia, University of Hull, University of Notre Dame, Penn State University, Columbia University, Stanford University, etc.*

2. La *Stanford University* tiene cerca de siete millones de registros y su tiempo de respuesta es de poco más de dos segundos.
<http://lib.stanford.edu/sdr>

3. Los tres primeros elementos son proyectos independientes que no están mantenidos por la propia comunidad.

4. Es el equivalente a una librería en Java o PHP, pero para Ruby.

5. Son los elementos principales de un objeto *Fedora*. Pueden ser xml o los propios ficheros digitales.

6. Es un estándar en la administración de proyectos desarrollado por el *Project Management Institute (PMI)*.
<http://www.pmi.org/PMBOK-Guide-and-Standards.aspx>

Bibliografía

Dushay, Naomi; Zumwalt, Matt (2013). *Hydra Stack - The hierarchy of promises*.
<https://wiki.duraspace.org/display/hydra/Hydra+Stack+-+The+Hierarchy+of+Promises>

EPrints (2010). *New features in EPrints 3.1*.
http://wiki.eprints.org/w/New_Features_in_EPrints_3.1

Green, Richard (2013). *Hydra objects, content models (cModels) and disseminators*.
<http://goo.gl/QGS0I>

Green, Richard; Zumwalt, Matt (2012). *Technical framework and its parts*.
<https://wiki.duraspace.org/display/hydra/Technical+Framework+and+its+Parts>

Fedora Project (2007). *The Fedora digital object model. Fedora release 3.0 Beta 1*.
<http://fedora-comons.org/documentation/3.0b1/userdocs/digitalobjects/objectModel.html>

The DSpace Foundation (2009). *DSpace system documentation: configuration and customization*.
http://www.dspace.org/1_5_2Documentation/ch05.html

Suscripción EPI sólo online

Pensando sobre todo en los posibles suscriptores latinoamericanos, ya no es obligatorio pagar la suscripción impresa de EPI para acceder a la online.

EPI se ofrece a instituciones en suscripción “sólo online” a un precio considerablemente más reducido (96,69 +21% IVA euros/año), puesto que en esta modalidad no hay que cubrir los gastos de imprenta ni de correo postal.