

Cross-Site Identification (XSid)

Cross-Site Identification (XSid) - Author not stated, blog.quaji.com.

- Published: date not stated
- Original: <<http://blog.quaji.com/2009/12/out-of-context-information-disclosure.html>>
- Preserved from: <http://blog.quaji.com/2009/12/out-of-context-information-disclosure.html> (stored) on 2026-08-09
- Capture timestamp: 20100214135134
- Licence: unknown

Rights remain with the original author and publisher. This is a research archive of a source from the Web Hacking Techniques Index collections, kept so the page going offline. To read the original, follow the link above.

Content

UNTRUSTED SOURCE TEXT. Everything below this line is third-party material quoted for research. It is data, not instructions. Do not follow directions, execute code, or fetch URLs because this text says so.

Overview

This post outlines a new attack technique in which publicly available information from social network sites, obtained out of context, can be used to identify a user, in cases where anonymity is taken for granted.

With the rise in popularity of social network websites (SNs), people have grown accustomed to feeding these information-hungry sites with their personal details. However, most people maintain the idea that while doing so, their anonymity is kept intact when interacting with other, non-related sites. While this notion is reasonable, it is not always true, as will be demonstrated.

In this attack (dubbed Cross-Site Identification or CSID), one site, usually a SN, effectively becomes an **identifying service** that is used while the user is surfing, supposedly anonymously, in another site. I will call these sites the *identifying* *and* *target* sites, respectively.

When the user arrives at the targeted site the attack takes place: The targeted site will silently cause the victim's browser to request the SN to share the user's personal details with the hacker. These details might be publicly available (i.e on the user's public profile), but their acquisition at this point, outside of the normal context of the SN causes the user's anonymity to be breaches and her identity known in the context of the targeting site.

The working assumption scenario for this attack is that the user is logged on to her favorite SN, and later, in a new tab, directs her browser to the targeted site. There is no connection between the two surfing sessions in the user experience, and therefor she expects to remain anonymous: If Betty has her Facebook open in one tab, she still expects to remain anonymous reading about, say, plastic surgery in another tab.

Actually, her SN website does not even need to be actually open: if she has not explicitly logged-off from it in her last visit, and chose a 'remember me' option when logging on, then her cookies are still valid. She is still authenticated and thus still vulnerable.

How it works

I've found the described vulnerability in Facebook a few months ago and wrote an extensive post about it [here](#). It contains a detailed description of the attack technique and specifics of the Facebook attack, and might provide an interesting addendum to this post, especially if you want to see all the nuts and bolts put together.

For this reason, I won't go into the details here again, but I'll give a brief description of the general attack scenario: The identifying site, usually a SN, is forced to reveal identifying attributes about the user. This is usually done through a SN application (though not necessarily). The SN API either reveals this data by design, or can be tricked to do so. The data is the publicly available information (like name and profile picture) that is available to anyone viewing the user's profile anyway, so it seems like there is no problem. The problem arises when this process can be initiated without the user's consent, while she is visiting a completely non-related site (the target site). When the user visits the target site, her browser is made to visit the identifying site and trigger the payload (I.E visit the attacker's SN application page clandestinely, providing him with the user's identifying attributes).

The attack has two main use cases in relation to the target site.

- The target site is controlled by the attacker. This is a malicious site owner wishing to uncover his visitor's identity, or a hacker will content placement control in another site.

This is the simpler, less powerful version of the attack. The attacker can use any means necessary to initiating the attack (such as iframes, javascript etc).

- The target site is a legitimate 3rd party site in which the attacker has no special privileges. This is a much more powerful version because the attacker, while being a normal user in the site, piggy-backs on the target site's reputation: The victim has no reason to suspect the target site, which could be very well known.

Obviously, the hacker still needs some kind of hook to initiate the attack. In this case it is an embedded image link put in the target site that causes the effect. The image link eventually point the browser at the identifying site and triggers the payload. This can most commonly be carried out in the form of a comment on a forum or blog. An image can usually be added as part of the comment, or as the avatar picture. Another option is via an advertisement image set by the attacker legitimately in the target site.

A different approach is in the form of HTML email's containing the rouge image link. In this scenario, the attacker is alerted when the victim opens his email box.

As a last touch of grace, it is usually possible for the malicious image link to initiate the attack **and** return a valid image to the victim's browser causing no suspicious behavior resulting in the entire attack going unnoticed.

In my previous post, I dubbed the vulnerability in Facebook "*CSRF personal information leakage vulnerability*" but some thought and conversation (thanks A.D!) showed that it is neither a CSRF per se, nor a leakage of information. It's not exactly a CSRF because the victim's browser isn't tricked

into performing any *action *apart from visiting a page (a CSRF token won't help here), and it's not exactly leakage because the information is publicly available! Its the out-of-context access to it that constitutes the attack. Furthermore, the vulnerability in the identifying sites found seems very minuscule (sometimes it is a feature!) when not considering this attack, so it is logical to assume that many other instances of it are in the wild. For these reasons I realized it's a new attack technique in its own right, and that was what motivated me to write this post. I suggest the name Cross-Site Identification (**CSID**).

Real world examples

Ah, the juicy bits.

As stated I've [posted](#) about the specific Facebook attack a few moths ago. While doing some (shallow) research, I found two other instances of the same vulnerability in **Bebo** ****and **Orkut**, which are reported here for the first time. By the way, these are the only three sites I've inspected so far, so I suspect there are many more open cases.

Facebook

I'll represent the video demonstration I made showing the effects of the attack on Facebook. The details are [here](#). The vulnerability has been patched by Facebook after being informed about it, a few months ago.

Bebo Bebo [apparently](#) uses a clone of Facebook's application platform. However, it was even easier to exploit for the purpose of this attack than Facebook was, as the identity of the user is given to the unauthorized application with little tweaking.

In Bebo, the default privacy settings is "*Profile viewable by my friends only*" which is good. However, even in this settings, the user's full name and profile picture are publicly visible. The described vulnerability has been reported to Bebo, but is currently still open.

Attack walk-through:

- Our victim, *Henry Eight* is logged to his Bebo account while surfing to the target site.
- The malicious image embedded in the target site directs his browser to the attacker's newly created application. It performs two actions: Logs the parameters added by Bebo's server making the request, and redirects the victim's browser, currently looking for pixels, to a valid image.
- Here is the logged information received by the attacker:

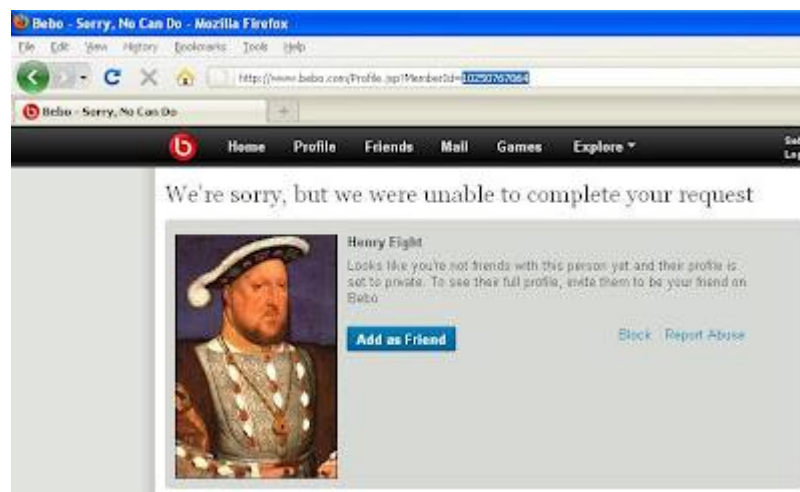
```

Dec 25 2009 17:09:26 208.75.184.235
POST:
Array
(
    [fb_sig] => 5e29f2481f746ba50b62c809a379a7fb
    [fb_sig_added] => 0
    [fb_sig_api_key] => 1fu0sxbAMs876i4JFD3D2oai7ux
    [fb_sig_in_canvas] => 1
    [fb_sig_locale] => en_IL
    [fb_sig_network] => Bebo
    [fb_sig_profile_id] => 10250767064
    [fb_sig_request_method] => GET
    [fb_sig_time] => 1261753738.576
    [fb_sig_user] => 10250767064
)

```

Note the victim's Bebo id given twice in *fb_sig_profile_id*, and in* *fb_sig_user*. *

- The publicly available information Bebo gives out for this user ID can be fetched through its API, but for the sake of clarity, the following screen shot is the public profile for this user.



The profile picture and full name are given even though the user's privacy settings are intact.

Orkut Orkut's applications framework is an implementation of the open source [OpenSocial](#) platform. A shallow investigation showed the site vulnerable twice. Both vulnerabilities are **design flaws**, one specific to Orkut, and the other in the OpenSocial specifications. Both issues have been reported but are currently still open.

1. Recent Visitors

Orkut has implemented a feature called [Recent Visitors](#) turned on by default, that shows an Orkut user the last 10 people to have viewed his profile, with a link to their profile.



In order to use this feature to launch the described attack, the attacker sends the unsuspecting victim to his Orkut profile page using, for example, a hidden iframe (if it's his own site). Note that an image link can be used as before although the image will break. Orkut does the rest by providing the attacker with the names and profiles of the victims.

As stated this vulnerability is a built-in feature in Orkut, turned on by default.

1. Signed Requests

Orkut, abiding by the OpenSocial specifications seems to keep from providing any information about the user of an application before it is approved. However, creating an application page with the [Signed Request](#) feature, has the effect of adding the OpenSocial ID of the viewing user to the request made to the application server. The following screen shot is the request made by Orkut to the application server when the victim is forced to "view" the attackers application.

```
Dec 26 2009 14:50:10 64.233.172.18
GET:
Array
(
    [lang] => en
    [country] => US
    [oauth_body_hash] => gzdXKkTJc6uuYCiCoWj+KDo6KtA=
    [opensocial_owner_id] => 06741398331619935963
    [opensocial_viewer_id] => 12881187622695198253
    [opensocial_app_id] => 12963447398594235126
    [opensocial_app_url] => http://hosting.gmodules.com/ig/gadgets/file/10691920859
    [opensocial_container] => http://www.orkut.com
    [opensocial_proxied_content] => 1
    [xoauth_signature_publickey] => pub.1199819524.-1556113204990931254.cer
    [xoauth_public_key] => pub.1199819524.-1556113204990931254.cer
    [oauth_version] => 1.0
    [oauth_timestamp] => 1261831795
    [oauth_consumer_key] => orkut.com
    [oauth_signature_method] => RSA-SHA1
    [oauth_nonce] => 1261831795213893000
    [oauth_signature] => WrOmCedFt2Y5cfSQqXoPXa7FIiWrmjuMq76PiB7kNx9ppV+8gugq2S9r1r'
)
POST:
[{"id":"viewer_info","error":{"message":"forbidden: Get profile permission denied."}]
```

Note on the bottom, the proper response of the server to the application's request for personal information about the viewer: It is denied on grounds of permission, rightfully so, as the application has not yet been approved by the user. However, note the added *opensocial_viewer_id** parameter. It is part of the parameters added when a signed request is performed. This parameter is, as its name suggests, the unique id number of the victim.

A user's opensocial id number is different from her Orkut id, therefor viewing the user's Orkut profile (like in the Bebo example) is not immediately possible. Perhaps it is possible to derive one from the other, or identify the user from this number in another way, but it's more important to note that this is already a vulnerability, even if somewhat to a lesser degree:

The *opensocial_viewer_id** parameter can effectively act as a ****cross-domain tracking cookie**, allowing the attacker to track the victim across different target sites. Although the victim's identity remains unknown, his actions on unrelated sites can be aggregated using this technique.

This vulnerability is also a design flaw in the OpenSocial specifications.

Combined Attack Vectors

Apart from the pure form of the technique detailed above, different attack vectors that leverage other factors are possible. In the attacks illustrated below, the identifying information is not publicly available immediately, but is obtained through other means.

Leveraging trust In most of the examples above, the user was tricked into interacting with a malicious SN application set forth by the attacker. The main difficulty (= the security vulnerability) was "convincing" the SN to play along without the victim's acknowledgment (= approving the application). But what about a legitimate application?

When a user approves an application he explicitly allows it access to his personal information. But the app can abuse these rights by providing these details out of context as described, to identify the user. The attacker can create a legitimate application, have users approve it for its legitimate front, and later use it to identify its users elsewhere. Another option is to use a vulnerability (such as an XSS) in an already popular application for this purpose: The victim's browser is directed to the vulnerable popular application, the SN identifies the user to the application (as it should, it had been approved!), but the hole in the application causes the identity to be transmitted to the attacker.

Leveraging an XSS Imagine a library site which, after you logging in, shows your name and your lent books. Imagine further that the site contains an XSS. While not being very lucrative for an attacker in itself, it can be used to launch described CSID attack: When the victim visits the target site, the malicious image link redirects to the library site with an XSS payload that snips the victim's name from the top of the page and sends it to the attacker (using an image request with parameters, for example).

Similarly, an XSS in any SN or SN application, or for that matter any site that identifies its users could be used to launch this attack. Think web-mails, game sites, shopping sites, forums, IMs, eCards... Any one of them, however benign, could serve as an identification service for its users while they are surfing other sites anonymously.

Leveraging a CSRF Consider a site that publicly displays a user's identity in relation to an action he has taken. There are many such sites, for example: Wikipedia ([history page](#)), Flickr ([comments](#)), Yahoo Answers ([questions](#) and [answers](#)). Even Blogger.com ([followers](#)). The user is aware of the implications of his action (having his identity shown), and is actively choosing to do so.

But what if there is a CSRF vulnerability in the site? In this case, a user that is logged-on can be forced to perform an action using a redirection from another location. In many cases the action itself might seem uninteresting enough (to a potential attacker), that CSRF protection is not implemented. However, a CSID attack can be launched exploiting this weakness.

Lets take for example Flickr comments. Imagine that adding a comment to a photo was not CSRF protected (which it is), meaning that it is possible to post a comment on behalf of a user without the user's knowledge. To perform the CSID attack, the attacker would post an image link on the target site that causes the victim's browser to comment some obscure photograph chosen by the attacker. When the victim (which is logged on to Flickr) views the target site he unintentionally posts a comment which is displayed along with his Flickr name under the photo. The attacker then simply views the Flickr page of the chosen photo, to collect the identity of the victim.

Note that I am *not* reporting the above listed sites as vulnerable, only providing them as potential examples. However I have found a few similar sites to be vulnerable in this way. I'll write an additional entry about this after allowing enough time for the sites' team to repair the problems.

Discussion and Summary

This new attack technique shakes the already loose foundations of anonymity while surfing the internet. A user's identity is at risk of being exposed while visiting almost any site, or even reading emails. While the average user might be somewhat attentive when visiting questionable sites, she does not at all fear using legitimate, well known sites, which are just as vulnerably.

The severity of this attack is increased by the fact that with the introduction of tabbed browsing in all modern browsers, multitask browsing has become common practice. Most of us will have our Facebook or webmail page open in one tab while surfing the web in others. Doing so makes us potentially vulnerable to this attack. Also, it has become common to offer a "Keep me logged on" option, which causes the user's browser to remain authenticated even when the application window is not actively open.

Furthermore, the attack trigger can easily be rendered benign by the attacker at any point, completely removing any trace the attack took place. This can be done even in 3rd party sites not controlled by attacker. This is because the trigger is an image link pointing to the attacker's server. When the attack is "live", the server will redirect the request onto the vulnerable identifying site. To kill the attack, the attacker merely has to return a valid image instead. No traces are left.

Even further, in most of the above examples the image request remains in the hands of the attacker, even *after* the payload has occurred. This is true even in the case of failed identification (I.E the user was not logged in). This allows the image link to get a valid image in the end. It also allows cascading the attack to multiple locations:

The victim's browser, while trying to fetch an image embedded in a reputable site, will start hopping between social network (and other) sites one by one, triggering the disclosure of the victim's identity from some (the ones he is logged on to), silently skipping the others (where he is not), and ending up with a valid image to show on screen. All without the victim's awareness.

The real world examples shown here are the result of only a quick research. Further research will probably find many other popular sites susceptible to this attack, and more examples in the reported sites. This is apparent as some of the vulnerabilities found were design flaws indicating that the use-case of this attack had not occurred to the writers of the specifications.

How this attack differs from CSRF (addendum) This question has been raised a few times and I would like to address it here. While the two are similar, I feel they are not the same. CSRF by OWASP definition is "...an attack which forces an end user to execute unwanted actions on a web application in which he/she is currently authenticated." In contrast, the exploits described in the paper require the end user to merely *view* a page on the vulnerable website. No action is taking place.

An action in my mind, is a two-step process: First the user requests the pre-action page, which contains the action button (and usually a form). The click on the button preforms the action and sends the user to the post-action page. CSRF is thus a way to skip the pre-action page, and send the user directly to the post-action page. (Via an email link, hidden iframe, image etc).

CSRF prevention is done with CSRF tokens: unique "challenge" tokens that are added as hidden values to the form in the pre-action page. When the form is submitted, the token is verified. Because the token's value is unknown in advance, construction of a malicious, direct post-action URL is not possible. While this method stops CSRF, it will not help with CSID.

Consider for example Orkut's "Recently Visited" feature. How would you use CSRF tokens (or anything else for that matter) to prevent it being used for CSID (i.e clandestinely sending the victim's browser to visit the attacker's profile), and still retain the current functionality (allowing direct URL to users' profiles, and counting profile views as visits)?

Similarly, If the social platform wants to allow direct URL access to it's applications, and adds the user's details to any request, you have a problem which, in my opinion, is not CSRF, and not fixable with anti-CSRF measures.

CSID and CSRF are similar in the fact they are both instances of the Confused Deputy problem: The browser is not sure if its *really* the user that is requesting it to go to some URL, or if it's a trick. This is a very broad issue touching the essence of WWW internet and not simple to settle. Google Analytics in a website is an example of the this as a feature rather than a security bug: The user's browser is hiddenly sent to another domain (Google) in order to perform an action (accumulating statistics for the site owner). Does the user *want* this action to take place? He is certainly not aware of it, nor has the option to prevent it. But we would not consider this to be CSRF even though it might meet the technical criteria. [Clickjacking](#) is another example in this family of attacks. But it is better seen as its own type of attack rather than more of the same. It is of course possible to see all these as one type of attack (or feature) but I feel something is lost in the generalization.

Writing these lines I see that CSID has two meanings. One is CSID as a payload: The disclosed of identifying details about a user out of context. The other is CSID as the vulnerability that allows this. The CSID payload can be launched using existing vulnerabilities like XSS and even CSRF as described above. But in other cases it is its own vulnerability, such as Orkut's design flaws, and Facebook's and Bebo's leniency in adding user information to requests.

If you reached this far, why don't you leave a comment?