

LSB 隐写术的密钥恢复方法^{*}

张卫明^{1,2}, 刘九芬^{1,2}, 李世取¹

(1. 信息工程大学信息工程学院, 河南 郑州 450002;
2. 中国科学院研究生院信息安全国家重点实验室, 北京 100039)

摘要: 作为隐写术安全性分析的重要手段, 隐写分析已成为信息隐藏领域的一个研究热点。而通过恢复隐写密钥来提取隐藏的消息是隐写分析的主要目的之一。为了研究如何搜索隐写密钥, 首先在“已知载体”和“载体被重复使用”条件下, 分析了恢复 LSB 隐写术密钥的计算复杂度。然后在“载体被重复使用”条件下, 对图像空域 LSB 隐写术提出了一种新的密钥恢复算法。该算法借鉴了密码分析中的“分别征服攻击”思想, 使计算复杂度由 $O(2^{2r})$ 降至 $O(2^r)$ 。实验结果表明了该算法的有效性。

关键词: 信息隐藏; 隐写分析; 提取攻击; 分别征服攻击

中图分类号: TP309 **文献标识码:** A **文章编号:** 0529-6579 (2005) 03-0029-05

隐写术和隐写分析是信息隐藏技术的重要分支。前者研究如何把秘密信息隐藏于可公开的多媒体数据中实现隐蔽通信, 后者研究如何检测、提取和破坏隐藏的秘密信息。

隐写算法包括嵌入算法和提取算法两部分。秘密信息通过嵌入算法被隐藏在载体对象中, 生成载密对象。接收者通过提取算法从载密对象中提取秘密信息。在秘密信息的嵌入和提取过程中通常会使用一个隐藏密钥。只有拥有此密钥的接收者才能检测或提取出隐藏的消息(遵循 Kerckhoffs 准则)。

隐写分析作为隐写术安全性分析的重要手段, 已成为信息隐藏领域的一个研究热点。目前在隐写分析领域的研究基本集中于隐蔽信息的检测^[1,2]。但是获取隐蔽信息往往是攻击者的终极目的, 我们把这种攻击称为“提取攻击”。而现在的隐写软件一般都要要求把消息先加密再隐藏, 即形成了“加密+隐藏”的安全通信模式, 当一个攻击者面对这种模式时, 他需要先做“提取攻击”, 然后才可以做传统的密码分析恢复明文信息。遗憾的是, 至今有关这方面研究的公开文献还很少。Chandramouli^[3]针对基于扩频通信的隐写术就一种特殊情况给出了提取攻击方法。他考察的情况是同一消息使用同一载体发送了两次, 前后两次的差别仅在于嵌入消息时所用的强度因子不同。在这种条件下, Chandramouli 使用盲源分离方法可以把消息从载密对象中分离出来。而对于目前普遍使用的对称密钥隐写术而言, 提取攻击只需恢复嵌入密钥即可。

Fridrich 等^[4]对基于图像 DCT 域(JPEG 图像)的 LSB 隐写算法(如 F5, Outguess), 利用卡方检验给出一种区分真伪隐写密钥的方法。

根据攻击者的已知条件可将“提取攻击”分为: “已知载体攻击”, “载体被重复使用的攻击”和“唯载密对象攻击”。众所周知, LSB 算法是隐写术中最简单也是最常见的算法。但这个算法却在隐写术中占有重要地位, 因为几乎全部的隐写算法中都可以找到 LSB 算法的影子, 互联网上常见的隐写软件中也大都使用 LSB 算法或 LSB 的衍生算法^[5]。由此, 本文研究一般 LSB 隐写术, 在“已知载体”和“载体被重复使用”条件下, 恢复隐写密钥的计算复杂度, 并在后一条件下, 对图像空域 LSB 隐写术提出一种密钥恢复方法, 同时给出了实验结果。

1 随机 LSB 隐写术模型

不失一般性, 我们以灰度级图像为例进行讨论。包含 N 个像素点的 256 色图像记作 N 维向量 $C^N = \{c_1, c_2, \dots, c_N\}$, $S^N = \{s_1, s_2, \dots, s_N\}$ 表示对应的载密图像, 其中 c_i 和 s_i 都是取值于 $[0, 255]$ 的整数, $1 \leq i \leq N$ 。用 $M^L = \{m_1, m_2, \dots, m_L\}$, $L \leq N$, 表示嵌入消息(一般为密文序列); $m_i \in \{0, 1\}$, $1 \leq i \leq L$ 。用 k 表示隐写密钥, 它取值于密钥空间 $\mathcal{K}$ 。

* 收稿日期: 2004-10-01

基金项目: 国家自然科学基金资助项目(60473022); 河南省自然科学基金资助项目(0511011300)

作者简介: 张卫明(1976年生), 男, 博士生; E-mail: nkdweiming@sohu.com

随机 LSB 隐写算法的嵌入过程如下: 首先利用隐写密钥 k 通过一个伪随机数发生器 G 生成随机序列 $y_1, y_2, \dots, y_L$, 然后按如下方式产生消息的随机嵌入位置 $x_i, 1 \leq i \leq L$

$$x_1 = y_1; x_i = x_{i-1} + y_i, 1 \leq i \leq L$$

最后把消息 $M = \{m_1, m_2, \dots, m_L\}$ 嵌入 $\{c_{x_1}, c_{x_2}, \dots, c_{x_L}\}$ 的 LSB 位, 从而得到载密图像 S , 具体的嵌入过程是: 若 C_{x_i} 的 LSB 位与 m_i 相同, 则不变, 否则进行“LSB 替换”(即把 C_{x_i} 的 LSB 位改成 m_i) 或按某种规则如通过“像素值 ± 1 ”实现消息嵌入, 合法的接收者拥有隐写密钥 k , 所以可以从载密图像 S 读出嵌入消息。

我们把随机数发生器的输出定义为独立同分布的随机变量序列 $Y_1, Y_2, \dots, Y_L$, 其取值为 $[a+1, a+d]$ 之间的整数, 其中 a 和 d 是两个整数, 满足

$$P\{Y_1 = a + i\} = p_i, 1 \leq i \leq d, \sum_{i=1}^d p_i = 1$$

易知只需讨论 $a=0$ 的情况即可 (若 $a \neq 0$, 可令 $Y'_i = Y_i - a$, 考察 Y'_i)。需要说明的是, 随机数发生器的输出一般要服从均匀分布。否则, 它在密码意义上是弱的, 因为在这种情况下密码分析者容易恢复其种子或构造等价的发生器, 所以我们重点分析 $p_i = \frac{1}{d}, 1 \leq i \leq d$ 的情况。此时最大随机间隔 d

由消息嵌入率 $\alpha = \frac{L}{N}$ 确定, 一般取 $d = \lceil \frac{2}{\alpha} - 1 \rceil$ 。

我们的目的是研究如何确定隐写密钥 k , 从而提取嵌入消息。

2 已知载体—单密钥碰撞攻击

如果攻击者除了接收到载密对象外, 还得到一部分载体 (比如同时拥有载密图像和部分载体图像)。此时攻击者通过比较, 可以观测到载体的部分被修改位置, 以及这些位置上的消息 (消息的一个随机抽样, 因为与消息相同的载体 LSB 位并没有被修改), 他的目的是通过这些信息恢复隐写密钥从而提取全部消息。一个自然的方法是穷举所有密钥, 生成相应的随机位置序列 $\{x_i\}$, 如果某个密钥生成的序列可以碰撞到所有观测到的修改位置, 则认为得到了真密钥, 否则作为伪密钥抛弃。我们把这种攻击称为“单密钥碰撞攻击”, 这里一个关键的问题是需要知道多少个修改位置才能确定出真密钥。

设 $Y_1, Y_2, \dots, Y_n, \dots$ 为相互独立且同分布

的随机变量序列, $P\{Y_i = i\} = \frac{1}{d}, 1 \leq i \leq d$ 。定义随机变量序列 $\{X_n, n \geq 1\}$

$$X_1 = Y_1; X_n = X_{n-1} + Y_n, n \geq 2$$

易知 $\{X_n, n \geq 1\}$ 为齐次马尔科夫链, 并有如下引理。

引理 1 对任给的 $n (n \geq 1)$ 个正整数 $x_1 < x_2 < \dots < x_n$, 下面的不等式成立 $P\{\exists \text{正整数 } i_1 < \dots < i_n, \text{ 使得 } X_{i_1} = x_1, \dots, X_{i_n} = x_n\} \leq \left[\frac{(d+1)^{d-1}}{d^d} \right]^n$

引理 1 保证了当密钥空间给定, 只要观测到的“载体被修改的位置”足够多就可以惟一确定隐写密钥。事实上, 由引理 1 易证得下面的定理。

定理 1 若隐写密钥的长度为 r 比特 (即 $|K| = 2^r$), 由上述单密钥碰撞攻击, 需要

$$n_0 = \left\lceil \frac{r}{d \log_2 d - (d-1) \log_2 (d+1)} \right\rceil$$

个“修改位置”可使伪密钥个数的数学期望小于 1, 从而可确定隐写密钥。

3 载体被重复使用条件下的攻击

3.1 双密钥碰撞攻击

我们讨论另一种情况: 即发送者从他的图像库中选择图像用以隐藏发送多组消息, 但同一图像被选择两次来发送两个 (或两段) 不同的消息。攻击者获得这样两幅载密图像后, 通过比较, 可记录二者的“差异位置”(即像素值不同的位置), 然后试验“每对”密钥, 如果某对密钥生成的随机位置可以覆盖所有“差异位置”, 则认为是真密钥对。我们把这种方法称为“双密钥碰撞攻击”。下面分析这种攻击的可行性, 并估计其需要的数据量。

设 $\{Y_{1n} \geq 1\}$ 和 $\{Y_{2n} \geq 1\}$ 是两条独立同分布的随机变量序列, 并且两条序列彼此相互独立 P

$\{Y_{j1} = i\} = \frac{1}{d}, 1 \leq i \leq d, j = 1, 2$, 定义两个随机变量序列 $\{X_{jnn} \geq 1\}, j = 1, 2$

$$X_{j1} = Y_{j1}, X_{jn} = X_{j,n-1} + Y_{jn}, n \geq 2$$

在引理 1 的基础上可证明下面结论。

引理 2 对任给的 $n (n \geq 1)$ 个正整数 $x_1 < x_2 < \dots < x_n$, 记 $p_n = P\{\exists 1 \leq s_1 < \dots < s_i \leq x_n, 1 < t_1 < \dots < t_j \leq x_n, \text{ 使得 } X_{1,s_1} = x_{1,1}, \dots, X_{1,s_i} = x_{1,i}, X_{2,t_1} = x_{2,1}, \dots, X_{2,t_j} = x_{2,j}, \text{ 且 } \{x_{1,1}, \dots, x_{1,j}\} \cup \{x_{2,1}, \dots, x_{2,j}\} = \{x_1, \dots, x_n\}\}$, 则

① 当 $d=2$ 时, $p_n \leq \left[\frac{121}{128} \right]^{\lceil \frac{n}{4} \rceil}$;

- ② 当 $d=3$ 时, $p_n \leq p_n \leq \left(\frac{80}{81}\right)^{\left\lceil \frac{n}{3} \right\rceil}$;
- ③ 当 $d \geq 4$ 时, $p_n \leq \left(\frac{2(d+1)}{d^d}\right)^n$

引理 2 说明当密钥空间给定, 只要已知的“差异位置”足够多就可以确定出一对真的隐写密钥。由引理 2 可得下面的定理。

定理 2 若隐写密钥的长度为 r 比特 (即 $|\mathcal{K}|=2^r$), 在已知重复使用载体的条件下, 需要 n_0 个“差异位置”可使伪密钥个数的数学期望小于 1, 其中 n_0 满足:

- ① 当 $d=2$ 时, $n_0 = \left\lceil \frac{8r}{7-2\log 11} \right\rceil - 4$;
- ② 当 $d=3$ 时, $n_0 = \left\lceil \frac{6r}{4\log 3 - \log 80} \right\rceil + 3$;
- ③ 当 $d \geq 4$ 时,
- $$n_0 = \left\lceil \frac{2r}{d\log_2 d - (d-1)\log_2 (d+1)} \right\rceil$$

3.2 分别征服攻击

上述方法虽然可以确定隐写密钥, 但是需要同时考虑两个密钥, 所以若密钥长度为 r 比特, 则攻击的计算复杂度为 $O(2^{2r})$ 。为了把复杂度降到 $O(2^r)$, 我们借鉴密码分析中“分别征服攻击”思想^[9], 把由两个密钥生成的两条位置序列当成一个序列密码的两条输入序列, 而把通过比较两幅载密图像得到的“差异位置”序列看作输出序列, 利用输入序列与输出序列的“相关性”分别恢复两个密钥。

设发送者重复使用的载体图像为 C^N , 把 C^N 的 LSB 位看作 0-1 随机变量序列 $B^N = \{B_1, \dots, B_N\}$, 即 $B_i, 1 \leq i \leq N$ 取值 $\{0, 1\}$ 。先后发送的两条消息记作 $M_1^L = \{M_{11}, \dots, M_{1L}\}$ 和 $M_2^L = \{M_{21}, \dots, M_{2L}\}$, 因为消息为密文, 所以可假设它们都是独立均匀分布的 0-1 随机变量序列, 两消息序列相互独立并都与 B^N 独立。设发送者嵌入这两条消息时使用的密钥为 k_1 和 k_2 , 由这两个密钥生成的嵌入位置分别记为 $X_1^L = \{X_{11}, \dots, X_{1L}\}$ 和 $X_2^L = \{X_{21}, \dots, X_{2L}\}$, 对应得到的两个载密图像记为 S_1^N 和 S_2^N , S_1^N 与 S_2^N 的差异位置记作 $D^q = \{d_1, \dots, d_q\}$ 。我们下面分析密钥 k_1 生成的前 n 个位置 $X_1^n = \{x_{11}, \dots, x_{1n}\}$ 中包含的差异位置的个数 ($U = |X_1^n \cap D^q|$) 的分布。任给 $1 \leq i \leq n$

(1) 若 $x_{1i} \notin X_1^n \cap X_2^n$ (即由密钥 k_2 没选到位置 x_{1i}), 则 $x_{1i} \in D_1^q$ 当且仅当 $M_{1i} \oplus B_{x_{1i}} = 1$ (即隐藏第一条消息时, 需对 x_{1i} 位置的像素值进行修改,

其中, $\oplus$ 表示模 2 加)。

(2) 若 $x_{1i} \in X_1^n \cap X_2^n$, 则存在 $1 \leq j \leq L$, 使得 $x_{2j} = x_{1i}$, 当 $M_{2j} \oplus B_{x_{1i}} = 0$ 时, $x_{1i} \in D_1^q$ 当且仅当 $M_{1i} \oplus B_{x_{1i}} = 1$ (即虽然隐藏两条消息时, 都选到了位置 x_{1i} , 但嵌第一条消息时需修改此位置的像素值, 而嵌第二条时不需要); 当 $M_{2j} \oplus B_{x_{1i}} = 1$ 时, $x_{1i} \in D_1^q$ 当且仅 $M_{1i} \oplus B_{x_{1i}} = 0$ (即嵌第二条消息时需修改 x_{1i} 位置的像素值, 而嵌第一条时不需要)。

构造 0-1 随机变量序列 $\{Z_i, 1 \leq i \leq n\}$: 当 $x_{1i} \notin X_1^n \cap X_2^n$, $Z_i = M_{1i} \oplus B_{x_{1i}}$; 当 $x_{1i} \in X_1^n \cap X_2^n$, $Z_i = (M_{2j} \oplus B_{x_{1i}} \oplus 1) \circ (M_{1i} \oplus B_{x_{1i}}) \oplus (M_{2j} \oplus B_{x_{1i}}) \circ (M_{1i} \oplus B_{x_{1i}} \oplus 1) = M_{1i} \oplus M_{2j}$, 即

$$Z_i = \begin{cases} M_{1i} \oplus B_{x_{1i}}, & \text{当 } x_{1i} \notin X_1^n \cap X_2^n \\ M_{1i} \oplus M_{2j}, & \text{当 } x_{1i} \in X_1^n \cap X_2^n \end{cases}$$

则 $U = X_1^n \cap D^q = \sum_{i=1}^n Z_i$ 。

由假设 X_1^L 和 X_2^L 都是独立均匀分布的 0-1 随机变量序列, 两序列相互独立并都与 B^N 独立, 易证得 $\{Z_i, 1 \leq i \leq n\}$ 也是独立均匀分布的 0-1 随机变量序列, 所以当充分大, 由中心极限定理知, 近似服从正态分布 $N(\frac{n}{2}, \frac{n}{4})$, 由强大数定律, U 以概率 1 接近 $\frac{n}{2}$ 。定义“ n 阶距离”为 $Dis(n) = |U - \frac{n}{2}|$, 此距离越小, 表示“输入”、“输出”间相关性越大, 故我们可根据“ n 阶距离”的大小来区分真伪密钥, 由此得到下面的攻击方法。

分别征服攻击算法:

现在已有许多检测方法可较精确的估计载密图像中的消息嵌入率 $\alpha^{[7]}$, 所以我们假设攻击者已知消息嵌入率, 从而他可估计出发送者使用的最大随机间隔 d , 然后攻击者进行如下各步操作:

(1) 比较两幅载密图像 S_1^n 与 S_2^n , 记录差异位置 $D^q = \{d_1, \dots, d_q\}$;

(2) 用定理 2 计算攻击所需的差异位置个数 n_0 。对于 $d \leq 3$, 令 $n_1 = n_0$; 对于 $d > 3$, 若 $(d+1)n_0 \leq q$, 令 $n_q = (d+1)n_0$, 否则, 令 $n_1 = q$, 取 $D^{n_0} = \{d_1, \dots, d_{n_0}\}$, $D^{n_1} = \{d_1, \dots, d_{n_0}\}$;

(3) 令 $n = \left\lceil \frac{2d_{n_1}}{d+1} \right\rceil$, 穷举密钥空间中的密钥,

用每个密钥 k 生成位置序列 $X_k^n = \{x_{k1}, \dots, x_{kn}\}$, 计算“ n 阶距离” $Dis_k(n) = \left| |X_k^n \cap D^{n_1}| - \frac{n}{2} \right|$, 把使

$Dis_k(n)$ 达到最小值或次小值的密钥存入备选集;

(4) 若 $|PSet|=2$, 则以 $PSet$ 中的一对密钥为真密钥, 用 3.3 节的方法做进一步的密钥区分; 若 $|PSet|>2$, 以 $PSet$ 为密钥空间, 对差异位置集 D^n 作“双密钥碰撞攻击”, 若能解出唯一一对密钥 $\{k_1, k_2\}$, 则以 $\{k_1, k_2\}$ 为真密钥做密钥区分, 否则攻击失败, 结束。

试验表明, 上述算法执行第 3 步后可满足 $|PSet|=2$, 即解出惟一的一对密钥 $\{k_1, k_2\}$, 此时算法的复杂度为 $O(2^r)$ 。易知, 此算法不依赖

于载体的性质。我们以 Lena.bmp 图象为例, 随机选择两个长度 $r=16$ 的密钥 k_1 和 k_2 , 用“Hide and Seek”^[8] 分别嵌入两条密文, 生成两幅载密图像, 然后做“分别征服攻击”, 只执行前 3 步, 取使 n 阶距离最小的两个密钥为真密钥。表 1 对 8 种不同的最大随机间隔 d 隔列出了实验数据, 其中当 $d=2$ 和 3 时, 使用的数据量很大, 这是因为定理 2 中对 $d=2$ 或 3 时所需数据量的估计比较粗糙, 实际上, 使用较少的数据也可使攻击成功。

表 1 分别征服攻击的试验数据

最大随机间隔 d (嵌入率)	2	3	4	5	6	7	8	9
	(0.67)	(0.50)	(0.40)	(0.33)	(0.29)	(0.25)	(0.22)	(0.20)
阶数 n	2390	7188	186	193	170	152	141	133
k_1 的 n 阶距离	31	5	3	1	4	7	5	1
k_2 的 n 阶距离	23	9	1	5	1	11	3	0
伪密钥最小 n 阶距离	82	770	7	20	19	16	14	11

3.3 密钥区分

用上述方法确定出真密钥对 $\{k_1, k_2\}$ 后, 我们还需要知道载密图像 S_1^N 与 S_2^N 分别是利用哪个密钥产生的。为此对两幅图像分别作均值滤波, 滤波值取实数, 得 S_1^N 和 S_2^N , 用密钥 k_1 生成位置序列 $X^L=\{x_1, \cdots, x_L\}$, 并以此分别对两幅载密图像抽样, 取使像素值为奇数的点, 计算

(1) $v_1=\frac{1}{L_1}\sum_{j=1}^{L_1}\left(s_{1,x_{ij}}-\bar{s}_{1,x_{ij}}\right)$ 其中 $x_{ij}\in X^L$ 并且 $s_{1,x_{ij}}$ 为奇数, $j=1, \cdots, L_1$;

(2) $v_2=\frac{1}{L_2}\sum_{h=1}^{L_2}\left(s_{2,x_{ih}}-\bar{s}_{2,x_{ih}}\right)$ 其中 $x_{ih}\in X^L$ 并且 $s_{2,x_{ih}}$ 为奇数, $h=1, \cdots, L_2$ 。

对于自然图像, 像素值与滤波值的差为图像的噪声部分, 近似服从 0 均值的正态分布。在“LSB 替换”模式下, 载密图像中所有像素值为奇数的点由两部分组成: 一部分与载体对应位置的像素值相同, 对于这种点, 像素值与滤波值的差可看成是来

自 0 均值正态分布的抽样; 而另一部分是由载体像素值加 1 得到的, 对于这种被修改的点, 其像素值与滤波值的差可看成是来自均值为 1 的正态分布的抽样。若 S_1^N 由 k_1 生成, 则 $\{s_{1,x_{ij}}, j=1, \cdots, L_1\}$ 中大约一半是由载体像素值加 1 得到的, 所以 $v_1\approx\frac{1}{2}$; 若 S_1^N 不由 k_1 生成, 而消息嵌入率为 α ($\alpha<1$), $\{s_{1,x_{ij}}, j=1, \cdots, L_1\}$ 中大约 $\frac{\alpha}{2}$ 是由载体像素值加 1 得到的, 所以 $v_1\approx\frac{\alpha}{2}$ 。由此我们可以做如下判断: 若 $v_1>v_2$, 则 S_1^N 由 k_1 生成, 否则 S_2^N 由 k_1 生成。

我们对 20 幅 8 比特灰度图像做了实验, 结果表明此方法可成功区分密钥。作为例子表 2 列出了以“Lena.bmp”和“Baboon.bmp”为载体时的实验结果。实验中, S_1^N 由密钥 k_1 生成, 所以 $v_1>v_2$ 表示密钥区分成功。

表 2 Lena.bmp 和 Baboon.bmp 图像密钥区分的试验数据

最大随机间隔 (嵌入率)	2	3	4	5	6	7	8	9
	(0.67)	(0.50)	(0.40)	(0.33)	(0.29)	(0.25)	(0.22)	(0.20)
Lena: $v_1(v_2)$	0.438 (0.357)	0.436 (0.346)	0.428 (0.322)	0.452 (0.315)	0.450 (0.249)	0.439 (0.274)	0.440 (0.247)	0.604 (0.306)
Baboon: $v_1(v_2)$	0.492 (0.384)	0.487 (0.334)	0.468 (0.321)	0.473 (0.337)	0.456 (0.311)	0.473 (0.293)	0.535 (0.319)	0.581 (0.357)

为叙述简便, 我们假设了两幅载密图像的消息嵌入率相同, 容易看出, 本节的方法同样适用于嵌入率不同的情况 (图 1)。

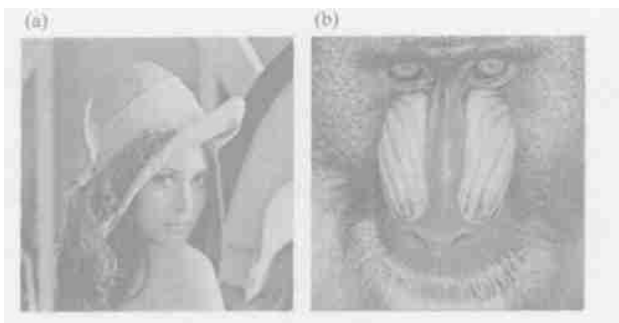


图 1 载体图像

Fig. 1 Cover images

4 结 论

本文主要研究 LSB 隐写术的密钥恢复问题。

首先分别分析了在“已知载体”条件下采用“单密钥碰撞攻击”、在“载体被重复使用”条件下采用“双密钥碰撞攻击”, 恢复隐写密钥的空间和时间复杂度。显然该分析既适用于空域又适用于频域, 既适用于图像又适用于音频和视频。

然后在“载体被重复使用”条件下, 借鉴密码分析中的“分别征服”思想, 对图像空域 LSB 隐写术提出了一种新的密钥恢复算法。该算法把时间复杂度由“双密钥碰撞攻击”的 $O(2^{2r})$ 降到了 $O(2^r)$, 试验结果表明该方法是有有效的。

需要说明的是, 由于本文是对一般的随机数发生器进行讨论的, 不涉及其结构和性质, 所以只能对密钥空间作简单的“穷举”。如果针对具体的随机数发生器, 则可结合密码分析中相应的分析方法, 进一步降低攻击的计算复杂度。

参考文献:

- [1] FRIDRICH J, GOLJAN M. Practical Steganalysis of Digital Images State of the Art // Security and Watermarking of Multimedia Contents of EI SPIE C]. 2002, 4675: 1—13.
- [2] ZHANG T, PING X J. A New Approach to Reliable Detection of LSB Steganography in Nature Images[J]. Signal Processing, 2003, 83(10): 2085—2094.
- [3] CHANDRAMOULI R. A Mathematical Framework for Active Steganalysis[J]. ACM Multimedia Systems Journal, Special Issue on Multimedia Watermarking, 2003, 9(3): 301—311.
- [4] FRIDRICH J, GOLJAN M, DU R. Searching for the Stego Key // Security, Steganography and Watermarking of Multimedia Contents of EI SPIE C]. 2004, 5306: 70—82.
- [5] JOHNSON N F. Steganography Tools. Available from: <http://www.jjtc.com/Security/stegtools.htm> 2005.
- [6] SIEGENTHALER T. Decrypting a class of stream ciphers using ciphertext only[J]. IEEE Transactions on Computers, 1985, C-34: 81—85.
- [7] FRIDRICH J, GOLJAN M. On Estimation of Secret Message Length in LSB Steganography in Spatial Domain // Security, Steganography and Watermarking of Multimedia Contents of EI SPIE C]. 2004, 5306: 23—34.
- [8] SHAGGY Hide, Seek. Available from: <http://www.jjtc.com/security/stegtools.htm>, 2005.

Approaches for Recovering Key of LSB Steganography

ZHANG Wei ming^{1,2}, LIU Jiu fen^{1,2}, LI Shi qu¹

(1. Institute of Information Engineering, Information Engineering University, Zhengzhou 450002, China;

2. State Key Laboratory of Information Security, Graduate School of the Chinese Academy Sciences, Beijing 100039, China)

Abstract: As an important way of security analysis for the steganography, steganalysis has become a concernful topic in the field of information hiding. And extracting the hidden message by recovering the stego key is one of main aims of steganalysis. To study how to search for the key of LSB steganography, firstly the theoretic analysis for the computational complexity of key recovery on LSB steganography are made under the condition of “known cover” and “cover being repeatedly used” respectively. Then under the latter condition, a key recovering method for LSB steganography of spatial images is presented, which can reduce the computational complexity from $O(2^{2r})$ to $O(2^r)$ by using the idea of “divide and conquer attack”. And the experiment results show that this method is effective.

Key words: information hiding, steganalysis, key recovery, divide and conquer attack