

AICR · OPEN PROBLEM

Asymptotically Optimal Gap Finders for Boolean-Terminal Martingales

The small-bias regime, where the Cleve–Impagliazzo bound degrades

Status. Statement: AI-written from Omid Etesami, Ji Gao, Saeed Mahloujifar and Mohammad Mahmoody, *Polynomial-time targeted attacks on coin tossing for any number of corruptions*, full version of a paper appearing in the Theory of Cryptography Conference (TCC) 2021, 2021, not yet checked by a human. Settled for $\mu = \Theta(1)$, where the Cleve–Impagliazzo gap finder already gives $\rho\alpha = \Omega(\mu^2/\sqrt{n}) = \Omega(\mu/\sqrt{n})$; open in the small- μ regime, where the known bound is a factor μ short, and the paper states it does not know the answer for general $\mu \leq 1/2$.

Category. *Information-Theoretic Cryptography.*

Conjecture (informal). *Think of a coin-tossing protocol in which n parties each broadcast one message and a Boolean output bit is then determined. Tracking the conditional probability that the output is 1 as the messages arrive gives a martingale that starts at the a-priori probability μ of the output being 1 and ends at 0 or 1. A classical theorem says such a martingale must jump noticeably somewhere: there is an online rule, which watches the messages one by one and must commit to a step without seeing the future, that stops at a step where the jump is large. The quality of such a rule is measured by two numbers, the probability that it stops at a large jump and the size of the jump it guarantees; their product is exactly what an adversary who may replace one message can convert into bias. When $\mu = \Theta(1)$ the classical rule already meets the conjectured target, since $\Omega(\mu^2/\sqrt{n}) = \Omega(\mu/\sqrt{n})$ there; no matching upper bound showing that rule optimal is proved anywhere. When μ is small, the classical rule degrades quadratically in it, and the best known guarantee for the product falls short of the target by a factor μ . The conjecture is that this loss is an artifact: a stopping rule achieving the target product should exist for every such martingale and every starting value at most one half. A second known family of rules controls only the expected jump rather than a probability-and-size pair, and the source paper does not know how to use that form in the targeted reduction it needs, so that route is unavailable to it; it does not show the form unusable in principle.*

1 The setting

Collective coin flipping [BOL90] asks how far an adversary who controls a few of n parties can push the output bit of a protocol in which each party broadcasts one message. The standard analytic handle is the Doob martingale of the output bit: if $b = f(w_1, \dots, w_n)$ is the Boolean output and $\mathbf{w}_i$ is the conditional probability that $b = 1$ given the first i messages, then $(\mathbf{w}_1, \dots, \mathbf{w}_n)$ is a martingale that starts at $\mu = \Pr[b = 1]$ and ends in $\{0, 1\}$.

Cleve and Impagliazzo [CI93] showed that such a martingale must jump: there is an *online* stopping rule — a “gap finder” — that, with probability $\Omega(\mu)$, stops at a step whose increment has magnitude $\Omega(\mu/\sqrt{n})$, i.e. a (ρ, α) gap finder with $\rho\alpha = \Omega(\mu^2/\sqrt{n})$. Their theorem was proved for $\mu = 1/2$; the same argument gives the

μ -dependent bound above, and it is exactly this dependence that degrades as μ shrinks. Khorasgani, Maji and Mukherjee [KMM19] proved a different guarantee, an *expected* gap of $2\mu(1 - \mu)/\sqrt{2n - 1}$, with no accompanying probability parameter; see also [KMW21].

Gap finders matter because they convert directly into attacks. It has long been known [CI93, KMM19] that a (ρ, α) gap finder yields a 1-replacing *non-targeted* attack of bias $\Omega(\rho\alpha)$: the adversary waits for the step where a large jump is about to occur, and either forces it (positive gap) or resamples to avoid it (negative gap). The source paper [EGMM21] closes the remaining half of that picture: it shows that a (ρ, α) gap finder also yields a *targeted* 1-replacing attack of bias $\Omega(\rho\alpha)$, one that pushes the output specifically towards 1, and that the reduction survives replacing the exact martingale by an approximation, hence runs in polynomial time. It also shows that a 1-replacing targeted attack composes recursively into a k -replacing one.

A positive answer would have let the paper use the gap-finder-derived targeted attack in place of its Lemma 47 and obtain information-theoretic k -replacing attacks by the recursive composition of its Section B (p. 45); the paper stresses that this route is included for completeness and that the results obtained this way are subsumed by its main theorem, Theorem 1 (p. 4). The paper achieves 1-replacing targeted bias $\Omega(\mu/\sqrt{n})$ directly, by a bespoke Azuma-based argument, so the target quantity is known to be achievable *as a bias*. But the route through gap finders currently loses a factor of μ , because the only available (ρ, α) gap finder has $\rho\alpha = \Omega(\mu^2/\sqrt{n})$, and the paper does not know how to use the expected-gap form of [KMM19] in the targeted reduction. The question is whether the gap-finding theorem itself can be strengthened to match.

2 Notation and parameters

- $n \in \mathbb{N}$ is the number of steps and $[n] = \{1, \dots, n\}$.
- $\mathbf{w}_{\leq n} \equiv (\mathbf{w}_1, \dots, \mathbf{w}_n)$ is a sequence of jointly distributed, discrete, real-valued random variables. For $i \in [n]$ we write $w_{\leq i} = (w_1, \dots, w_i)$ for a prefix in the support $\text{Supp}(\mathbf{w}_{\leq i})$. The sequence is a *martingale* if $\mathbb{E}[\mathbf{w}_{i+1} \mid \mathbf{w}_{\leq i} = w_{\leq i}] = w_i$ for

every $i \in [n - 1]$ and every $w_{\leq i} \in \text{Supp}(\mathbf{w}_{\leq i})$.

- The sequence is extended at both ends by the constant $\mathbf{w}_0 := \mathbb{E}[\mathbf{w}_1]$ and by $\mathbf{w}_{n+1} := \mathbf{w}_n$, so that the increments under consideration are $\mathbf{w}_i - \mathbf{w}_{i-1}$ for $i \in [n] \cup \{n + 1\}$, the last of which is identically 0.
- $\mu := \mathbb{E}[\mathbf{w}_1] = \mathbf{w}_0$ is the starting value of the martingale.
- r denotes the randomness of a stopping algorithm, $\tau \in [n] \cup \{n + 1\}$ its stopping time, and $|w_\tau - w_{\tau-1}|$ the *gap* it exhibits.
- $\rho \in [0, 1]$ and $\alpha > 0$ are the two parameters of a gap finder, and $c > 0$ is a universal constant.

The parameters of the statement are:

- n , the number of steps of the martingale (in the coin-tossing application, the number of parties, each sending one message).
- μ , the starting value of the martingale, $\mathbb{E}[\mathbf{w}_1]$; in the application, the probability that the protocol outputs 1 before any attack. Restricted to $(0, 1/2]$.
- α , the guaranteed size of the jump that the stopping rule exhibits at its stopping step.
- ρ , the probability, over the martingale and the rule's randomness, that the stopping step really does exhibit a jump of size at least α .
- τ , the step at which the online stopping rule stops; $n + 1$ if it never stops.
- c , a universal constant, independent of n , of μ , and of the martingale.

3 The conjecture

Definition 1 (Online stopping algorithm). Let $\mathbf{w}_{\leq n}$ be a sequence of jointly distributed discrete random variables. A (possibly randomized) algorithm *Stop* is an *online stopping algorithm* for $\mathbf{w}_{\leq n}$ if for every fixing r of its randomness:

1. $\text{Stop}_r(w_{\leq i}) \in \{0, 1\}$ for every $i \in [n]$ and every $w_{\leq i} \in \text{Supp}(\mathbf{w}_{\leq i})$, where the output 1 is read as “stop”; and
2. the decision is monotone: if $\text{Stop}_r(w_{\leq i}) = 1$ then $\text{Stop}_r(w_{\leq j}) = 1$ for every $j \geq i$ and every $w_{\leq j} \in \text{Supp}(\mathbf{w}_{\leq j})$ extending $w_{\leq i}$.

For $w_{\leq n} \in \text{Supp}(\mathbf{w}_{\leq n})$ the *stopping time* is

$$\text{StopTime}_r(w_{\leq n}) := \min\{i \in [n] : \text{Stop}_r(w_{\leq i}) = 1\}$$

if this set is non-empty, and $\text{StopTime}_r(w_{\leq n}) := n+1$ otherwise. The algorithm is not required to be efficient.

Definition 2 (((ρ, α) gap finder). Let $\mathbf{w}_{\leq n}$ be a martingale, extended by $\mathbf{w}_0 := \mathbb{E}[\mathbf{w}_1]$ and $\mathbf{w}_{n+1} := \mathbf{w}_n$. An online stopping algorithm *Stop* for $\mathbf{w}_{\leq n}$ is a (ρ, α) *gap finder* if, over the joint choice of $w_{\leq n} \leftarrow \mathbf{w}_{\leq n}$ and of the randomness r of *Stop*, writing $\tau = \text{StopTime}_r(w_{\leq n})$,

$$\Pr[|w_\tau - w_{\tau-1}| \geq \alpha] \geq \rho.$$

Note that stopping at $\tau = n+1$ contributes no gap, since $\mathbf{w}_{n+1} = \mathbf{w}_n$.

Conjecture 1 (Asymptotically optimal gap finders). *There is a universal constant $c > 0$ — independent of n , of μ , and of the martingale — such that the following holds. For every $n \in \mathbb{N}$, every $\mu \in (0, 1/2]$, and every martingale $\mathbf{w}_{\leq n} \equiv (\mathbf{w}_1, \dots, \mathbf{w}_n)$ with $\mathbb{E}[\mathbf{w}_1] = \mu$ and $\Pr[\mathbf{w}_n \in \{0, 1\}] = 1$, extended by $\mathbf{w}_0 := \mu$ and $\mathbf{w}_{n+1} := \mathbf{w}_n$, there exist a real $\alpha > 0$ and an online stopping algorithm*

Stop for $\mathbf{w}_{\leq n}$ such that Stop is a (ρ, α) gap finder for $\mathbf{w}_{\leq n}$ with

$$\rho \cdot \alpha \geq \frac{c\mu}{\sqrt{n}},$$

where

$$\rho = \Pr_{w_{\leq n} \leftarrow \mathbf{w}_{\leq n}, r} [|w_\tau - w_{\tau-1}| \geq \alpha]$$

and $\tau = \text{StopTime}_r(w_{\leq n})$. No bound is imposed on the running time of Stop; the statement is information-theoretic.

Remark 1 (What is known). The paper records the two known guarantees side by side (its Theorem 52): a $(\Omega(\mu), \Omega(\mu/\sqrt{n}))$ gap finder from Cleve–Impagliazzo, and an expected-gap bound $\alpha = 2\mu(1 - \mu)/\sqrt{2n - 1}$ from Khorasgani–Maji–Mukherjee. It notes the first was proved for $\mu = 1/2$ and that the same proof carries over to general μ with the stated μ -dependence. It proves the converse half of the motivation — that a (ρ, α) gap finder yields a targeted 1-replacing attack of bias $\Omega(\rho\alpha)$ (its Theorem 53), robust to ε -approximate martingale access (its Theorem 60) — and it separately achieves targeted 1-replacing bias $\Omega(\mu/\sqrt{n})$ by a direct argument (its Lemma 47), which shows the target product is not larger than what an attack can already extract by other means.

Remark 2 (Openness). The paper’s own formulation is: “Then, can one always obtain an (ρ, α) gap finder for $\mathbf{w}_{\leq n}$ such that $\rho\alpha = \Omega(\mu/\sqrt{n})$?” (p. 45), introduced by “This motivates the following question, which as far as we know is open” (p. 45). The reason it is asked: “Unfortunately, the bound of the (ρ, α) gap finder of [CI93] degrades with μ , which means we cannot use their result to obtain an asymptotically similar result to that of our Lemma 47” (p. 45), together with “We do not know how to find (expected) α gap finders for this purpose. This means we cannot use the result of [KMM19] (see Theorem 52)” (p. 45). On the consequence of a positive answer: “A positive answer to the following question above would have allowed us to use the derived targeted attack instead of our Lemma 47 and obtain (information-theoretic) k -replacing attacks through recursive compositing, as stated in Section B” (p. 45). On the settled regime, the paper notes of the Cleve–Impagliazzo bound that “their result is only proved for almost unbiased final

bits in which both $\{0,1\}$ happen with probability $\Omega(1)$ ” (p. 44).

Remark 3 (Caveats for a reviewer). Four things to check. First, the printed question omits the hypothesis that the martingale’s final value lies in $\{0,1\}$; taken literally it is false, since a martingale constant at μ has no gaps at all. That hypothesis has been restored above, which the surrounding text forces — the appendix opens by describing the Cleve–Impagliazzo theorem as being about “any n -step martingale in which the final value is in $\{0,1\}$ ”, and the reduction it feeds uses the Doob martingale of a Boolean function. A reviewer should confirm this is the right repair and not, say, a $[0,1]$ -boundedness hypothesis instead. Second, the heading calls the target “asymptotically optimal”, which suggests a matching $\rho\alpha = O(\mu/\sqrt{n})$ upper bound, but no such upper bound was found proved or cited anywhere in the paper; Conjecture 1 as stated does not depend on it. Third, the printed question has been read as existential in the pair (ρ, α) , uniformly over n , μ and the martingale; a reader could instead read ρ or α as pinned down in advance. Fourth, the Khorasgani–Maji line of work on estimating martingale gaps was active around and after 2021; whether a later paper in that line settles this has not been checked, and it is the first place to look.

Remark 4 (Why it is worth settling). This is a quantitative strengthening of a thirty-year-old structural theorem about bounded martingales, in the one regime where that theorem is known to be lossy, and the loss is exactly quadratic in the parameter. A positive answer would give a clean modular derivation of targeted coin-tossing attacks that are optimal up to constants for every corruption budget: gap finder, then the paper’s targeted reduction, then its recursive composition. A negative answer would be at least as interesting, since it would separate what online stopping rules can certify about a martingale from what a one-message-replacing adversary can actually achieve against the same protocol — the paper’s own Lemma 47 already reaches bias $\Omega(\mu/\sqrt{n})$, so a refutation would show the gap-finder abstraction is strictly weaker than the attacks it was introduced to explain. The statement itself mentions no protocols, no adversaries and no cryptography: it is a question about real-valued martingales with Boolean terminal value and online stopping rules, stateable in half a page from Definitions 1 and 2, with one inequality, one universal constant, and two explicit bounds

to beat.

4 Bibliography

- [BOL90] Michael Ben-Or and Nathan Linial. *Collective coin flipping*. Advances in Computing Research, 5:91–115, 1990.
- [CI93] Richard Cleve and Russell Impagliazzo. *Martingales, collective coin flipping and discrete control processes*. Manuscript, 1993.
- [EGMM21] Omid Etesami, Ji Gao, Saeed Mahloujifar, and Mohammad Mahmoody. *Polynomial-time targeted attacks on coin tossing for any number of corruptions*. Theory of Cryptography Conference (TCC); full version, 2021. [UNVERIFIED]
- [KMM19] Hamidreza Amini Khorasgani, Hemanta K. Maji, and Tamalika Mukherjee. *Estimating gaps in martingales and applications to coin-tossing: constructions and hardness*. Theory of Cryptography Conference (TCC), pages 333–355, Springer, 2019.
- [KMW21] Hamidreza Amini Khorasgani, Hemanta K. Maji, and Mingyuan Wang. *Optimally-secure coin-tossing against a byzantine adversary*. 2021 IEEE International Symposium on Information Theory (ISIT), pages 2858–2863, 2021.