

Constant-Overhead Boolean AMD Circuits

Size $O(|C|) + \text{poly}(\lambda) \cdot |C|^{0.9}$, resisting every additive attack on the wires

Status. Statement: AI-written from Elette Boyle, Geoffroy Couteau, Niv Gilboa, Yuval Ishai, Lisa Kohl, Nicolas Resch and Peter Scholl, *Oblivious Transfer with Constant Computational Overhead*, IACR ePrint 2023/817. Not yet checked by a human. Polylogarithmic overhead is achieved by Genkin, Ishai and Weiss; constant overhead is open, and the source’s Theorem 43 shows that achieving it would — together with the source’s own constant-overhead OT protocol — settle the main open question on constant-overhead secure computation for general Boolean circuits.

Category. *Secure Multiparty Computation.*

Conjecture (informal). *An algebraic manipulation detection circuit computes what it is supposed to compute while making tampering useless: whatever an attacker achieves by toggling any subset of the internal wires, it could have achieved by toggling bits of the input and output instead. Such circuits are what turn a semi-honest protocol into a maliciously secure one, and their size overhead is what that security costs. The best known construction pays a polylogarithmic factor. The source proves that a constant factor would be enough to settle one of the main open questions about the asymptotic complexity of cryptography — securely computing Boolean circuits against malicious parties with only constant computational overhead — because it supplies the other missing ingredient, a constant-overhead protocol for oblivious transfer.*

1 The setting

Securely computing Boolean circuits with constant computational overhead is, as the source puts it, “[o]ne of the main open questions about the asymptotic complexity of cryptography”. In the semi-honest model it is settled: local PRGs suffice, by Ishai, Kushilevitz, Ostrovsky and Sahai. Against malicious parties it was posed as an open question in that same work, and the best known protocols pay a factor growing polylogarithmically with the security parameter and the circuit size.

The source’s main result is the first constant-overhead protocol for malicious bit-OT. That does not immediately give general functionalities: OT is complete for secure computation, but the best known protocols for Boolean circuits in the OT-hybrid model have polylogarithmic overhead, whereas in the *semi-honest* OT-hybrid model the textbook GMW protocol achieves perfect security at small constant overhead. So the gap between semi-honest and malicious survives the source’s result — and what the source does with it is reduce the general question to a question about circuits.

The reduction runs through algebraic manipulation detection (AMD) circuits, introduced by Genkin, Ishai, Prabhakaran, Sahai and Tromer. An AMD circuit computes f while resisting additive attacks: the effect of toggling any subset of its wires can be simulated by an ideal additive attack on the inputs and outputs of f alone. Applying the semi-honest GMW protocol to an AMD circuit, with a suitable encoding protecting input and output, yields a maliciously secure protocol; combined with a constant-overhead OT protocol this reduces the main open question to the design of constant-overhead AMD circuits.

The state of the art is Genkin, Ishai and Weiss’s construction with polylogarithmic overhead in $|C|$ and λ . The source records: “Whether this can be improved was left open, but the question was further reduced to the design of two kinds of simple protocols in the honest-majority setting: a protocol that only provides semi-honest security (with a constant fraction of corrupted parties) and a protocol that only guarantees the correctness of the output.” It contrasts this with the route of Ishai, Prabhakaran and Sahai and of Damgård, Ishai and Krøigaard, which instead reduces the question to honest-majority protocols with malicious security.

2 Notation and parameters

- $C : \{0,1\}^n \rightarrow \{0,1\}^k$ is a deterministic or randomized Boolean circuit; $|C|$ is its size, and the overhead of an implementation is measured against $|C|$.
- $\widehat{C}$ is the AMD implementation, itself a randomized Boolean circuit with the same input and output lengths.
- An *additive attack* A toggles an arbitrary chosen subset of the wires of $\widehat{C}$; $\widetilde{C} \leftarrow A(\widehat{C})$ is the tampered circuit.

- Δ_{in} over $\{0,1\}^n$ and Δ_{out} over $\{0,1\}^k$ are the ideal additive attacks the simulator is allowed on the input and the output.
- ϵ is the security error in statistical distance; λ is the security parameter, and $\epsilon = 2^{-\lambda}$ is the setting Theorem 43 needs.
- $\text{SD}(\cdot, \cdot)$ is statistical distance.

3 The conjecture

Definition 1 (Boolean AMD circuit, [BCGIKRS23, Definition 42], after [GIP+14]). Let $C : \{0,1\}^n \rightarrow \{0,1\}^k$ be a (deterministic or randomized) Boolean circuit. A randomized Boolean circuit $\widehat{C} : \{0,1\}^n \rightarrow \{0,1\}^k$ is an ϵ -secure AMD implementation of C if

- *Completeness*: for all $x \in \{0,1\}^n$, $\widehat{C}(x) \equiv C(x)$; and
- *Security against additive attacks*: for any additive attack A , toggling a subset of the wires of $\widehat{C}$, there exist distributions Δ_{in} over $\{0,1\}^n$ and Δ_{out} over $\{0,1\}^k$ such that for every $x \in \{0,1\}^n$,

$$\text{SD}\left(\widehat{C}(x), C(x \oplus \Delta_{\text{in}}) \oplus \Delta_{\text{out}}\right) \leq \epsilon,$$

where $\widetilde{C} \leftarrow A(\widehat{C})$.

Conjecture 1 (Constant-overhead AMD circuits). *Every Boolean circuit C admits a $2^{-\lambda}$ -secure AMD implementation $\widehat{C}$ (Definition 1) of size*

$$O(|C|) + \text{poly}(\lambda) \cdot |C|^{0.9}.$$

Theorem 1 (What it buys, [BCGIKRS23, Theorem 43], cf. [GIW16, Claim 18]). *Suppose Conjecture 1 holds. Then a constant-overhead OT protocol implies a constant-overhead protocol for general Boolean circuits.*

Remark 1 (Why the exponent 0.9, and why the shape matters). The size bound is not $O(|C|)$ outright: it allows an additive $\text{poly}(\lambda) \cdot |C|^{0.9}$ term. That is what makes the statement the right one to conjecture — the second term is sublinear in $|C|$, so for circuits

large relative to the security parameter the total is $O(|C|)$, while the slack absorbs per-circuit machinery that need not be linear. The exponent 0.9 is the source's, transcribed as printed; any constant strictly below 1 would serve the same purpose, and a resolution achieving a different exponent should say so rather than being read as settling this exact form.

Remark 2 (The other half is already done). Theorem 1 is a conditional statement with two hypotheses, and the source discharges one of them: its main result is a constant-overhead protocol for malicious bit-OT. So Conjecture 1 is the whole of what stands between the literature and a constant-overhead maliciously secure protocol for general Boolean circuits. That is a stronger position than the question was in before the source appeared, and it is why the conjecture is worth stating on its own.

Remark 3 (Both directions, and what a negative answer would mean). The source poses this as a design question and predicts nothing. A refutation — a proof that AMD circuits must pay more than constant overhead — would not settle the main open question negatively, because Theorem 1 is only one direction: constant-overhead protocols might exist without going through AMD circuits, via the honest-majority route of Ishai, Prabhakaran and Sahai. So a lower bound here closes a route, not the problem.

Remark 4 (Partial results the source already has, which are not this). The source obtains constant-overhead protocols for general functionalities under *relaxed* security in its Section 6.1, and constant-overhead protocols for some restricted classes of functionalities. Neither settles Conjecture 1, and neither is what this statement asks: the conjecture is about AMD circuits for every Boolean circuit at full $2^{-\lambda}$ security.

4 Bibliography

[BCGIKRS23] Elette Boyle, Geoffroy Couteau, Niv Gilboa, Yuval Ishai, Lisa Kohl, Nicolas Resch and Peter Scholl. *Oblivious transfer with constant computational overhead*. IACR ePrint 2023/817. The source. Definition 42 and Theorem 43 are on pages 32–33, and the statement of the main open question is in Section 6, page 32.

- [IKOSo8] Yuval Ishai, Eyal Kushilevitz, Rafail Ostrovsky and Amit Sahai. *Cryptography with constant computational overhead*. 40th ACM Symposium on Theory of Computing (STOC), 2008, pages 433–442. Settles the semi-honest case from local PRGs, and poses the malicious case as the open question the source’s reduction addresses.
- [GIP+14] Daniel Genkin, Yuval Ishai, Manoj M. Prabhakaran, Amit Sahai and Eran Tromer. *Circuits resilient to additive attacks with applications to secure computation*. 46th ACM Symposium on Theory of Computing (STOC), 2014, pages 495–504. Introduces AMD circuits, and the source’s Definition 42 is theirs.
- [GIW16] Daniel Genkin, Yuval Ishai and Mor Weiss. *Binary AMD circuits from secure multiparty computation*. Theory of Cryptography Conference (TCC) 2016-B, Part I, LNCS 9985, pages 336–366. The state of the art: AMD circuits with polylogarithmic overhead, and the source of the claim Theorem 43 cites.
- [DIK10] Ivan Damgård, Yuval Ishai and Mikkel Kroigaard. *Perfectly secure multiparty computation and the computational overhead of cryptography*. EUROCRYPT 2010, LNCS 6110, pages 445–465. The polylogarithmic-overhead protocols the source’s question is measured against.
- [IPSo8] Yuval Ishai, Manoj Prabhakaran and Amit Sahai. *Founding cryptography on oblivious transfer — efficiently*. CRYPTO 2008, LNCS 5157, pages 572–591. The alternative route, which reduces the question to honest-majority protocols with malicious security rather than to AMD circuits.
- [GMW87] Oded Goldreich, Silvio Micali and Avi Wigderson. *How to play any mental game or a completeness theorem for protocols with honest majority*. 19th Annual ACM Symposium on Theory of Computing (STOC), 1987, pages 218–229. The semi-honest protocol that, applied to an AMD circuit, yields the malicious one.