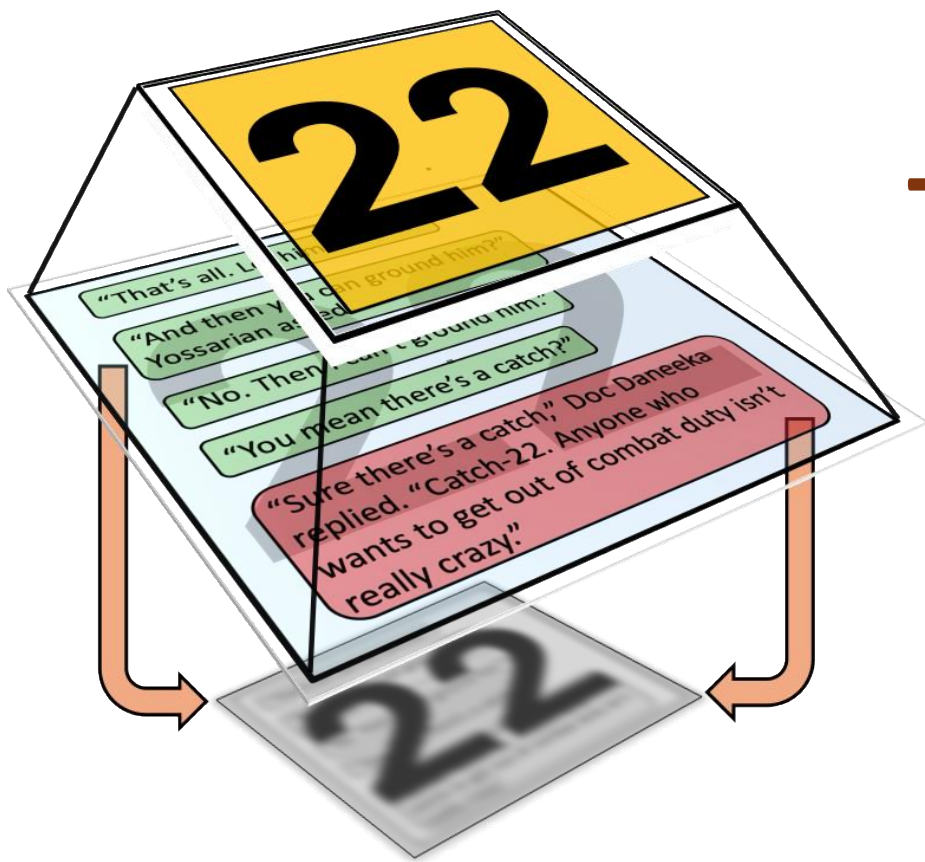




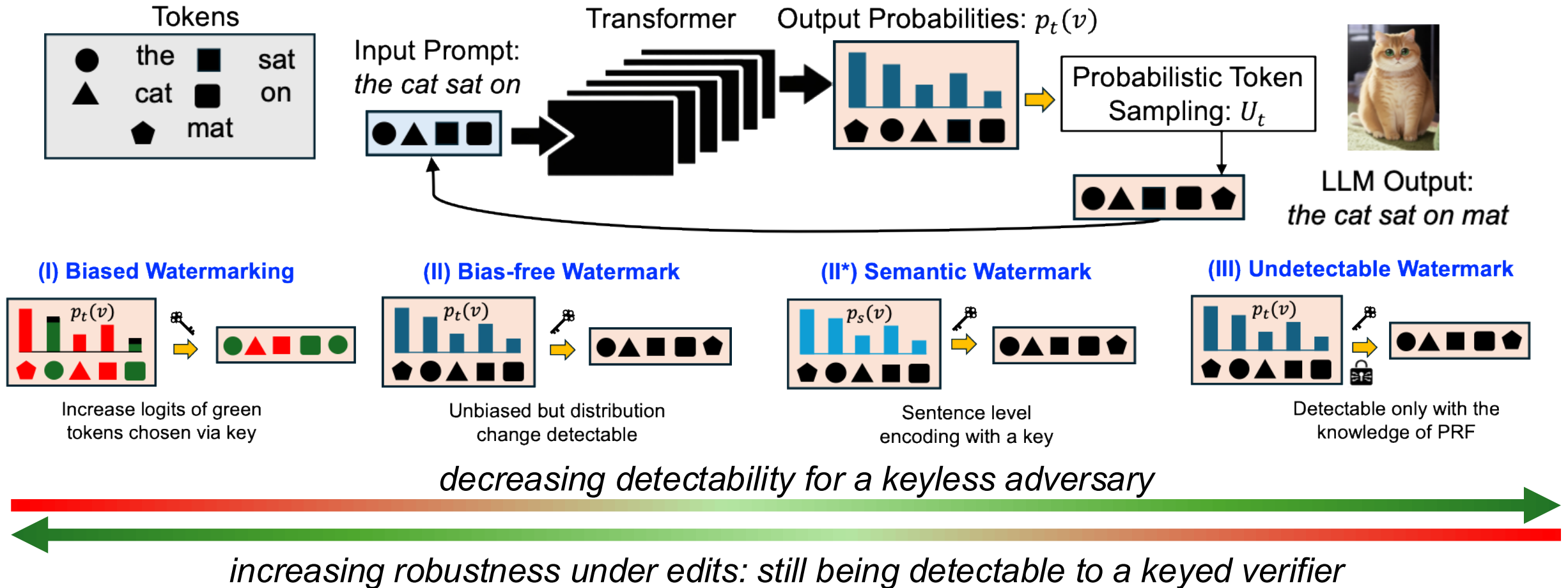
Catch-22: On the Fundamental Tradeoff Between Detectability and Robustness in LLM Watermarking

Kuheli Pratihari and Debdeep Mukhopadhyay,
Indian Institute of Technology Kharagpur, India



ICML
International Conference
On Machine Learning

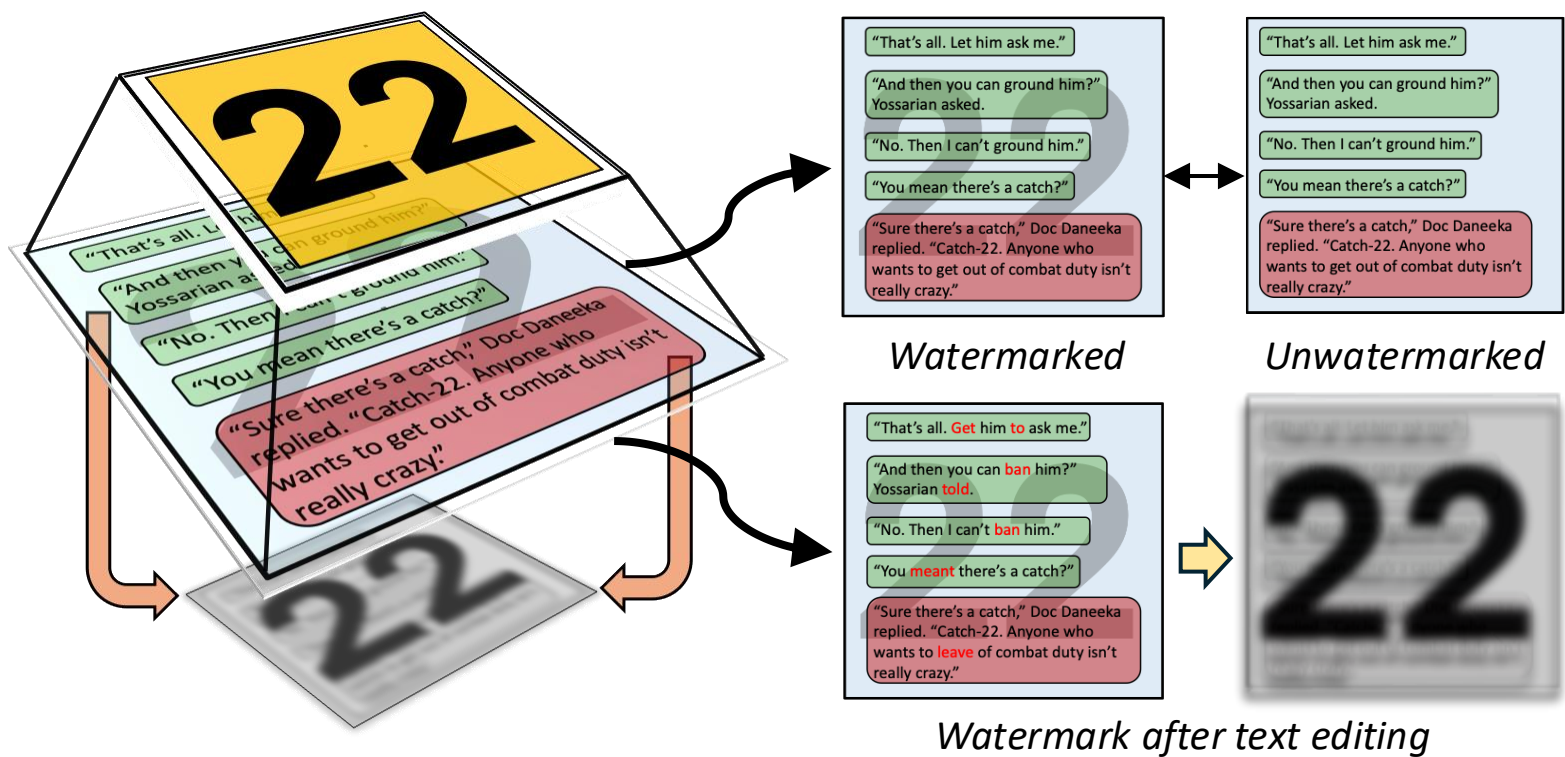
Detectability-Robustness Tradeoff



Since 2023 [Kirchenbauer et al.] watermark design research has developed significantly, but:

- (i) lack a unified theoretical standpoint that can compare different watermark families
- (ii) tie a scheme's robustness and keyless detectability directly to construction

Quantifying detectability & robustness across different watermarks



Sampling family (T output tokens)	Detectability (TV bound)
Greedy	$O(1)$
Biased (δ tilt)	$O(\delta \sqrt{T})$
Bias-free	$O(\sqrt{T})$
Semantic (avg. sentence length l)	$O(\sqrt{T/l})$
Distribution- preserving	0

We quantify the upper bound on detectability in terms of total variation distance

Sampling family (T output tokens)	Post-edit KL budget	Maximum Tolerable edit-rate
Token-level (biased / bias-free)	$T_s(1 - \epsilon)^2 D_0^{(\text{tok})}$	$1 - \sqrt{\log_2(1/\beta)/TD_0^{(\text{tok})}}$
Semantic	$T_s(1 - 2\epsilon_s)^2 D_0^{(\text{sem})}$	$\leq \frac{1}{2} \left(1 - \sqrt{\log_2(1/\beta)/TD_0^{(\text{sem})}} \right)$

Robustness of watermark after edits to the LLM output decay quadratically

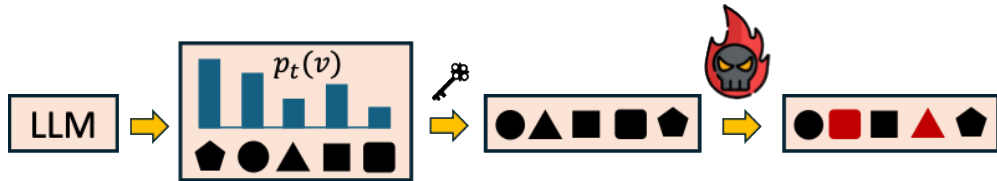
Theorem 1 [informal]: Any probability-modifying watermark leaves a statistical trace whose strength grows with the length of the generated text.

Theorem 2 [informal]: Watermark detection power contracts quadratically under text edits, imposing a fundamental trade-off between stealth and robustness.

Theoretical Contributions: Theorem 1 and 2 are formally defined in paper as Theorem 3.2 and 3.4.

Experiments across different watermarking schemes

Theorem 3 [informal]: The hybrid selection rule provably picks the watermark scheme with the smallest stealth cost that still meets the target robustness under the anticipated edit regime.



Edit rate estimation $\rightarrow \epsilon \rightarrow$ hybrid selection rule

$$S(m) = \text{AUC}(m, \hat{\epsilon}) + 1/(1 + \max(z(m, \hat{\epsilon}), 0))$$

Score is computed for a target LLM at the estimated edit rate.

Token-based watermarks including biased & bias-free ($\hat{\epsilon} \approx 0.25$, $\hat{\epsilon}_s \approx 0.06$)		Semantic-based Watermarks ($\hat{\epsilon} \approx 0.42$, $\hat{\epsilon}_s \approx 0.1$)	Distribution preserving ($\hat{\epsilon} \approx 0$, $\hat{\epsilon}_s \approx 0$)
★ HCW: 1.137 (0.910, 3.40)	DiPMark: 1.104 (0.90, 3.90)	★ Pmark: 1.305 (0.850, 1.20)	★ CGW: 1.990 (0.99, -5.80)
HeavyWater: 1.072 (0.880, 4.20)	SimplexWater: 1.052 (0.870, 4.50)	SemStamp: 1.220 (0.820, 1.50)	
Unigram: 0.982 (0.880, 8.80)	KGW: 0.954 (0.86, 9.60)	SimMark: 1.170 (0.800, 1.70)	

Tuple for each watermark $\langle \text{Watermark} \rangle: S(m, \hat{\epsilon})$ $(\text{AUC}(m, \hat{\epsilon}), z(m, \hat{\epsilon}))$	HCW chosen for ($\hat{\epsilon} \approx 0.25$, $\hat{\epsilon}_s \approx 0.06$) PMark chosen for ($\hat{\epsilon} \approx 0.42$, $\hat{\epsilon}_s \approx 0.1$) CGW chosen for ($\hat{\epsilon} \approx 0$, $\hat{\epsilon}_s \approx 0$)
---	---

We empirically find that our proposed hybrid rule attains the highest score across 8 watermarking families on Llama-2 and Mistral-7B models

Thank you for visting!



You are cordially invited to my
Poster (**ID 66807**)



Wed, Jul 8, 2026



from 5:00 PM – 6:45 PM KST



**COEX Convention & Exhibition Center,
HALL A, Seoul, South Korea**

We are thankful for the partial funding from the Center for Hardware-Security Entrepreneurship Research & Development (C-HERD) and Information Security Education and Awareness (ISEA), both initiatives of MeitY, Govt. of India.