



POLITECNICO
MILANO 1863

A NOVEL SELF-NORMALIZED BERNSTEIN-LIKE DIMENSION-FREE INEQUALITY AND REGRET BOUNDS FOR GENERALIZED KERNELIZED BANDITS

ALBERTO MARIA METELLI, SIMONE DRAGO, AND MARCO MUSSI

{albertomaria.metelli, simone.drago, marco.mussi}@polimi.it



MOTIVATION

Linear Bandits (LBs, Abbasi-Yadkori et al., 2011): expected reward linear in an unknown d -dimensional parameter vector θ^* :

$$\mathbb{E}[y_t | \mathbf{x}_t; \theta^*] = f^*(\mathbf{x}_t) = \langle \mathbf{x}_t, \theta^* \rangle$$

Generalized Linear Bandits (GLBs, Filippi et al., 2010): apply a known *inverse link function* μ :

$$\mathbb{E}[y_t | \mathbf{x}_t; \theta^*] = \mu(f^*(\mathbf{x}_t)) = \mu(\langle \mathbf{x}_t, \theta^* \rangle)$$

- ✓ Different noise models (exponential family)
- ✗ Finite-dimensional $d < \infty$ **only**

Kernelized Bandits (KBs, Chowdhury and Gopalan, 2017): unknown function f^* in *reproducing kernel Hilbert space* (RKHS) with known kernel k :

$$\mathbb{E}[y_t | \mathbf{x}_t; \theta^*] = f^*(\mathbf{x}_t), \quad f^*(\mathbf{x}) = \langle f^*, k(\mathbf{x}, \cdot) \rangle$$

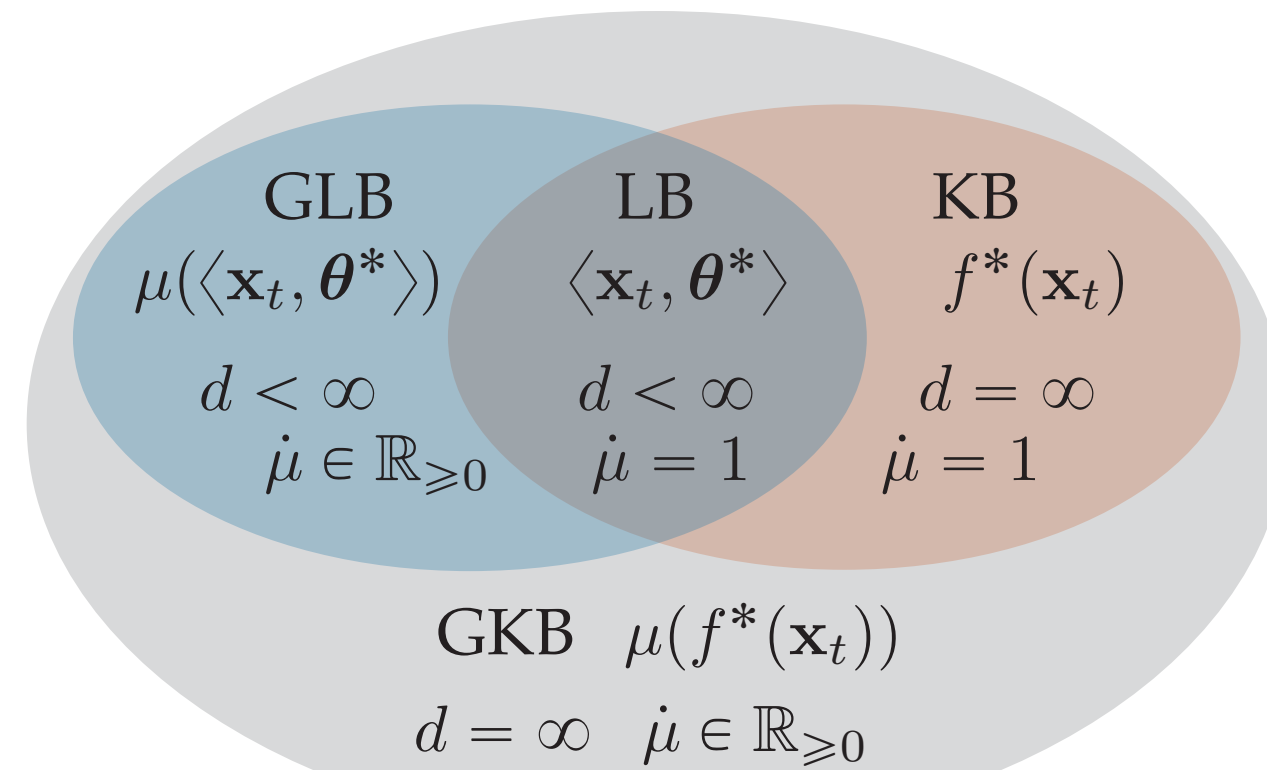
- ✗ Subgaussian noise model **only**
- ✓ Infinite-dimensional $d = \infty$

SETTING

Generalized Kernelized Bandits (GKBs, ours): unknown function f^* in RKHS with known kernel k + apply a known *inverse link function* μ :

$$\mathbb{E}[y_t | \mathbf{x}_t; \theta^*] = \mu(f^*(\mathbf{x}_t)), \quad f^*(\mathbf{x}) = \langle f^*, k(\mathbf{x}, \cdot) \rangle$$

$$\text{GKB} := \text{GLB} + \text{KB}$$



PRELIMINARIES

CANONICAL EXPONENTIAL FAMILY

$$p(y | \mathbf{x}; f) = \exp \left(\frac{yf(\mathbf{x}) - m(f(\mathbf{x}))}{g(\tau)} + h(y, \tau) \right)$$

- m is 3 times differentiable and convex
- **Inverse link function**: $\mu = m'$

$$\mathbb{E}[y | \mathbf{x}; f] = \mu(f(\mathbf{x})), \quad \text{Var}[y | \mathbf{x}; f] = \frac{\dot{\mu}(f(\mathbf{x}))}{g(\tau)}$$

RKHS AND INFORMATION GAIN

- Kernel k induces a countable-dimensional feature mapping ϕ s.t. $f(\mathbf{x}) = \langle \alpha, \phi(\mathbf{x}) \rangle$

- **Weighted kernel and feature**:

$$\tilde{\phi}(\mathbf{x}; f) = \sqrt{g(\tau)^{-1} \dot{\mu}(f(\mathbf{x}))} \phi(\mathbf{x})$$

$$\tilde{k}(\mathbf{x}, \mathbf{x}'; f) = g(\tau)^{-1} \sqrt{\dot{\mu}(f(\mathbf{x}))} k(\mathbf{x}, \mathbf{x}') \sqrt{\dot{\mu}(f(\mathbf{x}'))}$$

- **Weighted info gain**: $\frac{1}{2} \log \det(\lambda^{-1} \tilde{\mathbf{K}}_t(\lambda; f))$

$$\tilde{\mathbf{K}}_t(\lambda; f) = \lambda \mathbf{I} + (\tilde{k}(\mathbf{x}_i, \mathbf{x}_j; f))_{i,j \in [t-1]}$$

- **Weighted covariance operator**:

$$\tilde{V}_t(\lambda; f) = \sum_{s=1}^{t-1} \tilde{\phi}(\mathbf{x}_s; f) \tilde{\phi}(\mathbf{x}_s; f)^\top + \lambda \mathbf{I}$$

$$\det(\lambda^{-1} \tilde{V}_t(\lambda; f)) = \det(\lambda^{-1} \tilde{\mathbf{K}}_t(\lambda; f))$$

ASSUMPTIONS

ASSUMPTIONS ON RKHS

f^* belongs to RKHS $\mathcal{H}$ with known kernel k :

- Asm 3.1 (**Bounded norm**): $\|f^*\| \leq B$
- Asm 3.2 (**Bounded kernel**): $\sup_{\mathbf{x} \in \mathcal{X}} k(\mathbf{x}, \mathbf{x}) \leq K$

ASSUMPTIONS ON THE REWARD

Reward y_t sampled as $y_t \sim p(\cdot | \mathbf{x}; f)$:

- Asm 3.3 (**Bounded noise**): $|\epsilon_t| := |y_t - \mu(f^*(\mathbf{x}))| \leq R$
- Asm 3.4 (**Self-concordance**, Russac et al., 2021): $|\ddot{\mu}(f(\mathbf{x}))| \leq R_s \dot{\mu}(f(\mathbf{x}))$

SELF-NORMALIZED BERNSTEIN-LIKE DIMENSION-FREE CONCENTRATION INEQUALITY

We need a concentration bound:

- **Bernstein-like**: handle effectively the **heteroskedasticity** of the noise ϵ_t
- **Dimension-free**: no dependence on **dimensionality** of features ϕ (can be infinite)
- **Empirical**: dependence on the **weighted covariance operator** $\tilde{V}_t(\lambda; f^*)$ (is random)

Self-normalized Concentrations

Dani et al. (2008)
Abbasi-Yadkori et al. (2011)
Chowdhury and Gopalan (2017)
Fauray et al. (2020)
Zhou et al. (2021)
Ziemann (2025)

Condition

Hoeffding
Hoeffding
Hoeffding
Bernstein
Bernstein
Bernstein

Dim-free

✗
✓
✓
✗
✗
✗

Properties

Empirical
✓
✓
✓
✗
✓

Heterosc.

✗
✗
✗
✓
✗
✗

Technique

Freedman
Pseudo-Max
Pseudo-Max
Pseudo-Max
Freedman
PAC-Bayes

Our work

Bernstein

✓

✓

✓

Freedman

$$\forall t \geq 1: \|S_t\|_{\tilde{V}_t^{-1}(\lambda)} \leq \hat{B}_t(\delta; f) := \left(\sqrt{73 \log \det(\lambda^{-1} \tilde{V}_t(\lambda))} + \sqrt{3} \right) \sqrt{\log \frac{\pi^2(\rho + 1)^2}{3\delta} + \frac{3RK}{\sqrt{\lambda}} \log \frac{\pi^2(\rho + 1)^2}{3\delta}} \text{ w.p. } 1 - \delta$$

$$\text{where: } S_t := \sum_{s=1}^{t-1} \epsilon_s \phi(\mathbf{x}_s), \quad \tilde{V}_t(\lambda) := \sum_{s=1}^{t-1} \sigma_s^2 \phi(\mathbf{x}_s) \phi(\mathbf{x}_s)^\top + \lambda \mathbf{I}, \quad \rho = \max \left\{ 0, \left\lceil \log \left(\frac{8R^2 K^2 (t-1)^3}{\lambda} \log \left(1 + \frac{K^2 R^2}{\lambda} \right) \right) \right\rceil \right\}$$

ALGORITHM: GKB-UCB

ALGORITHM

For every round $t \in [T]$:

1. **Maximum Likelihood Estimate**:

$$\hat{f}_t \in \arg \min_{f \in \mathcal{H}} \mathcal{L}_t(f) := \sum_{s=1}^{t-1} \frac{-y_s f(\mathbf{x}_s) + m(f(\mathbf{x}_s))}{g(\tau)} + \frac{\lambda}{2} \|f\|^2$$

2. **Optimistic Decision Selection**: $(\tilde{f}_t, \mathbf{x}_t) \in \arg \max_{f \in \mathcal{C}_t(\delta), \mathbf{x} \in \mathcal{X}} \mu(f(\mathbf{x}))$

$$\text{where: } g_t(f) := \sum_{s=1}^{t-1} \frac{\mu(f(\mathbf{x}_s))}{g(\tau)} \phi(\mathbf{x}_s) + \lambda \alpha$$

$$\mathcal{C}_t(\delta) = \left\{ f \in \mathcal{H} : \|g_t(f) - g_t(\hat{f}_t)\|_{\tilde{V}_t^{-1}(\lambda; f)} \leq \sqrt{\lambda} B + g(\tau)^{-1} \hat{B}_t(\delta; f) \right\}$$

REGRET ANALYSIS

If $\lambda \geq \Omega(K^2)$, w.p. $1 - \delta$: $R(\text{GKB-UCB}, T) := \sum_{t \in [T]} (\mu(f^*(\mathbf{x}^*)) - \mu(f^*(\mathbf{x}_t))) \leq$

$$\tilde{O} \left((1 + R_s B K) \left(\sqrt{\lambda} B + \sqrt{\tilde{\gamma}_T(\mathcal{H}) \log \frac{1}{\delta}} + R K \log \frac{1}{\delta} \right) \sqrt{\tilde{\gamma}_T(f^*)} \sqrt{\frac{T}{\kappa_*}} \right)$$

Maximal weighted information gains:

$$\tilde{\gamma}_t(f) := \max_{\mathbf{x}_1, \dots, \mathbf{x}_{t-1} \in \mathcal{X}} \frac{1}{2} \log \det(\lambda^{-1} \tilde{\mathbf{K}}_t(\lambda; f))$$

$$\tilde{\gamma}_t(\mathcal{H}) := \sup_{f \in \mathcal{H}} \tilde{\gamma}_t(f)$$

Complexity index:

$$\kappa_* := \frac{1}{\dot{\mu}(f^*(\mathbf{x}^*))}$$

REFERENCES

- Y. Abbasi-Yadkori, D. Pál, and C. Szepesvári. Improved algorithms for linear stochastic bandits. In *Advances in Neural Information Processing Systems*, 2011.
- S. R. Chowdhury and A. Gopalan. On kernelized multi-armed bandits. In *International Conference on Machine Learning*, 2017.
- V. Dani, T. P. Hayes, and S. M. Kakade. Stochastic linear optimization under bandit feedback. In *Conference on Learning Theory*, 2008.
- L. Fauray, M. Abeille, C. Calauzènes, and O. Fercoq. Improved optimistic algorithms for logistic bandits. In *International Conference on Machine Learning*, 2020.
- S. Filippi, O. Cappé, A. Garivier, and C. Szepesvári. Parametric bandits: The generalized linear case. In *Advances in Neural Information Processing Systems*, 2010.
- Y. Russac, L. Fauray, O. Cappé, and A. Garivier. Self-concordant analysis of generalized linear bandits with forgetting. In *International Conference on Artificial Intelligence and Statistics*, 2021.
- D. Zhou, Q. Gu, and C. Szepesvári. Nearly minimax optimal reinforcement learning for linear mixture markov decision processes. In *Conference on Learning Theory*, 2021.
- I. Ziemann. A vector bernstein inequality for self-normalized martingales. *Transactions on Machine Learning Research*, 2025.

PAPER

Take a look
at our work!

