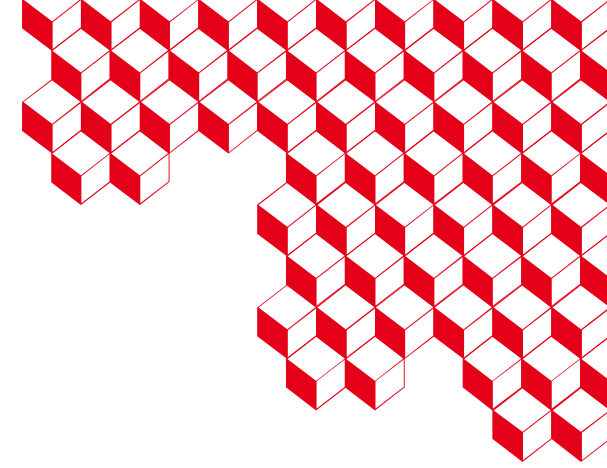




Active subspace embedding for parsimonious gaussian process surrogates

Amandine MARREL*, Bertrand IOOSS‡

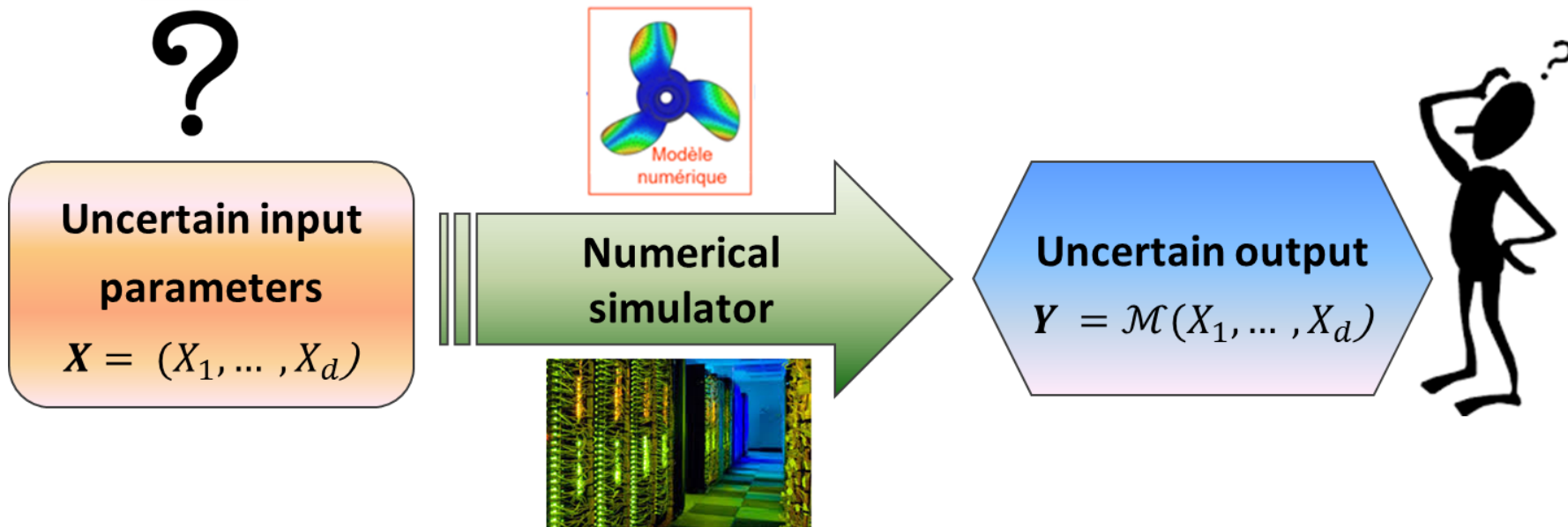
* CEA Energy Division, IRESNE, DER, Cadarache, France

‡ EDF R&D, Chatou, France

ENBIS Conference 2026, Florence, Italy.

Risk assessment in nuclear accident analysis

- **Safety studies:** compute a failure risk (margins, rare events) with validated computer/numerical models
- **Numerical simulators:** fundamental tools to understand, model & predict physical phenomena
- **Large number of input parameters**, related to physical and numerical modelling
- **Uncertainty on some inputs** → **uncertainty on output & safety margins**



Risk assessment in nuclear accident analysis

- How to deal with uncertainties in numerical simulation?

- Probabilistic framework and Monte Carlo-based methods

- **CPU-expensive simulator** $\Rightarrow$ Use of a **surrogate model*** to propagate input uncertainties

- **Applicative constraints/framework:**

- ✓ **Given data:** a single i.i.d. inputs/output **sample** $(x^{(i)}, y_i)_{1 \leq i \leq n}$ where $y_i = \mathcal{M}(x^{(i)}) \rightarrow$ to be used for multi-purpose (*sensitivity analysis, uncertainty propagation... And training a metamodel*)

- $\rightarrow n$ -dataset denoted: $\mathcal{D}_n = (\mathbf{X}_s, \mathbf{Y}_s)$

- ✓ **Small sample size:** $n \approx 100$ to 1000 simulations

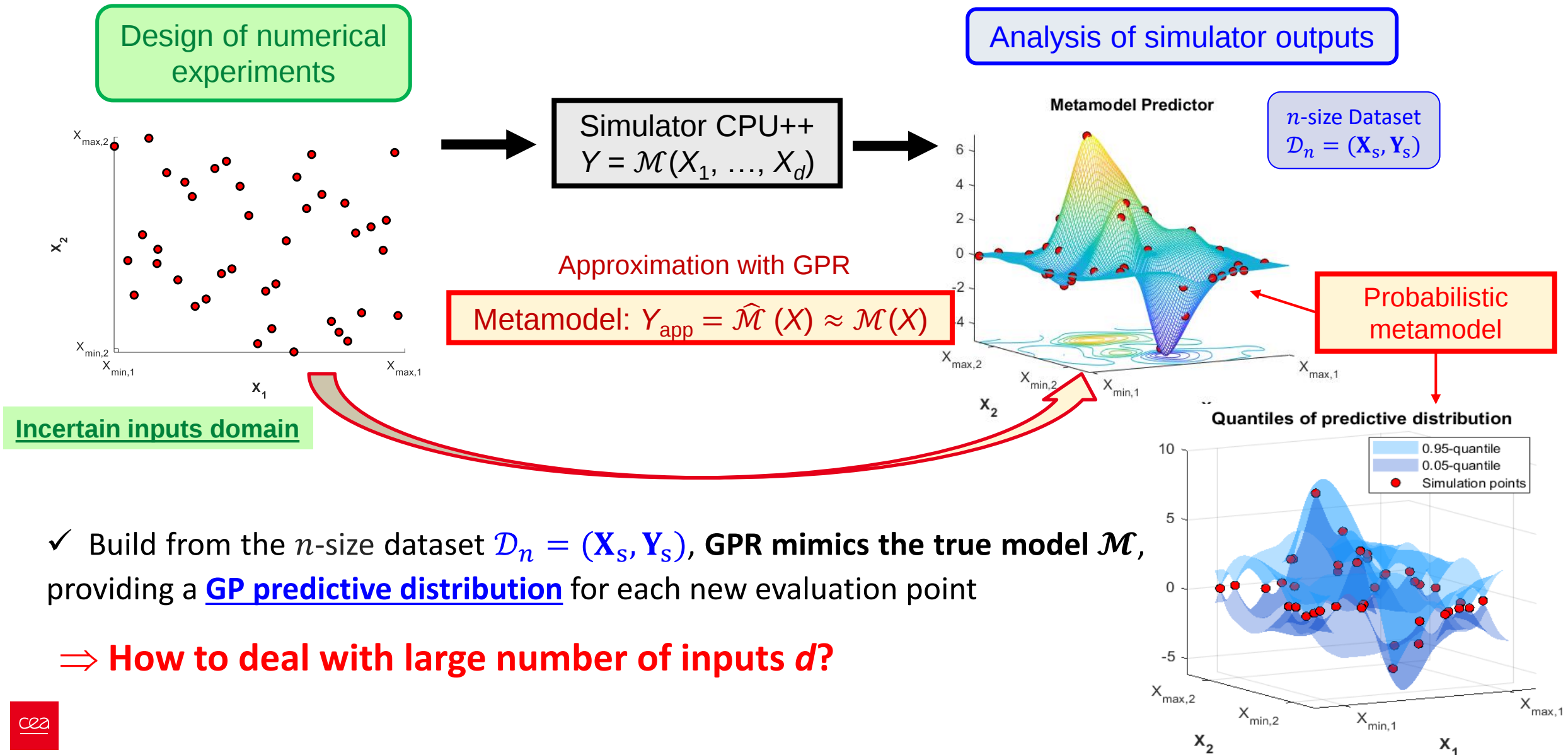
- ✓ **Large number of uncertain inputs:** $d \approx 10$ to 100 inputs

- ✓ **Required predictive uncertainty associated to each metamodel prediction**



Gaussian Process Regression (GPR): particularly well-suited tool $\Rightarrow$ Very popular

Crucial use of GPR metamodel



How to build parsimonious GPR in large dimension?

1. Reminds on Gaussian Process Regression

Reminders on GPR

► **Probabilistic surrogate model**: response considered as a realization of a random GP [RW06,Gra21]

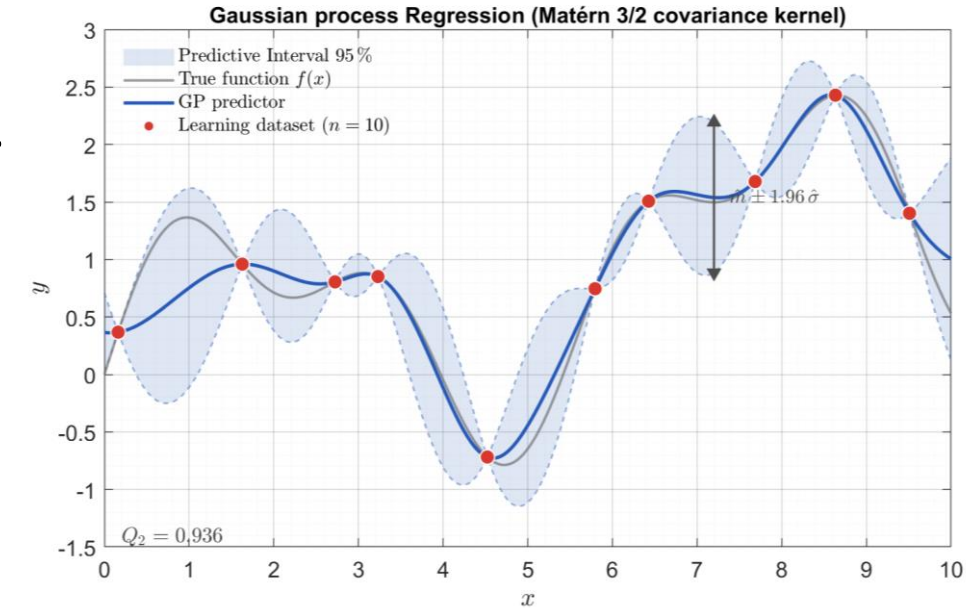
$$Y(\mathbf{x}) \sim GP(\mu(\mathbf{x}), k(\mathbf{x}', \mathbf{x}))$$

With $\mu(\mathbf{x})$ the mean and $k(\mathbf{x}', \mathbf{x})$ the covariance function.

⇒ **Predictive GP** is the GP conditioned by the dataset $\mathcal{D}_n$:

$$Y(\mathbf{x}^*)|_{\mathcal{D}_n} \sim GP(\hat{y}(\mathbf{x}^*), \hat{s}(\mathbf{x}', \mathbf{x}^*))$$

With analytical formulations for $\hat{y}(\mathbf{x}^*)$ and $\hat{s}(\mathbf{x}', \mathbf{x}^*)$



⇒ **Conditional mean** $\hat{y}(\mathbf{x}^*)$ serves as the **predictor** at location $\mathbf{x}^*$

⇒ **Conditional covariance** $\hat{s}(\mathbf{x}^*, \mathbf{x}^*)$ gives the **predictive variance** at location $\mathbf{x}^*$

⇒ **Predictive interval** of any **level** α can be built at any location $\mathbf{x}^*$

Recommendations for parametric choices

► **In practice:** standard parametric choices for trend function μ and covariance function k

$$Y(\mathbf{x}) \sim GP(\mu(\mathbf{x}), k(\mathbf{x}', \mathbf{x}))$$

⇒ For μ : either **constant** or linear basis

⇒ For k : stationary covariance built-upon tensorized 1-D covariance functions of ν -Matérn

$$\text{1-Dim} \rightarrow k_{\sigma, \nu, \theta}(x, \tilde{x}) = \sigma^2 \frac{2^{1-\nu}}{\Gamma(\nu)} \left(\frac{\sqrt{2\nu}h}{\theta} \right)^\nu K_\nu \left(\frac{\sqrt{2\nu}h}{\theta} \right)$$

3/2 or 5/2 Matérn covariances
offer good properties and
« intermediate » regularity

$$\text{d-Dim} \rightarrow k_{\sigma, \nu, \theta}(\mathbf{x}, \tilde{\mathbf{x}}) = \sigma^2 \prod_{i=1}^d k_{1, \nu, \theta_i}(x_i - \tilde{x}_i) \text{ with } h = |\mathbf{x} - \tilde{\mathbf{x}}| \Rightarrow \text{Hyperparameters } \boldsymbol{\theta} \in \mathbb{R}^{+,d}$$

	$\nu = \frac{1}{2}$	$\nu = \frac{3}{2}$	$\nu = \frac{5}{2}$	$\nu = +\infty$
Usual name	exponential	3/2-Matérn	5/2-Matérn	Gaussian
$k_{\sigma, \nu, \theta}(x, \tilde{x})$	$\sigma^2 e^{-\frac{h}{\theta}}$	$\sigma^2 (1 + \sqrt{3}\frac{h}{\theta}) e^{-\sqrt{3}\frac{h}{\theta}}$	$\sigma^2 \left(1 + \sqrt{5}\frac{h}{\theta} + \frac{5}{3} \left(\frac{h}{\theta} \right)^2 \right) e^{-\sqrt{5}\frac{h}{\theta}}$	$\sigma^2 e^{-\frac{1}{2} \left(\frac{h}{\theta} \right)^2}$
Differentiability of GP trajectories	$\mathcal{C}^0$	$\mathcal{C}^1$	$\mathcal{C}^2$	$\mathcal{C}^\infty$

⇒ Additional nugget effect : $\tilde{k}(\mathbf{x}', \mathbf{x}) = k(\mathbf{x}', \mathbf{x}) + \tau^2 \delta_{\mathbf{x}=\mathbf{x}'}$ (→ nugget hyperparameter τ)

Estimation of GP hyperparameters (θ, τ)

⇒ How to robustly estimate the **hyperparameters** $(\theta, \tau) \in \mathbb{R}^{+,d+1}$ from the learning sample?

➔ For good **predictivity** + **reliable prediction intervals** ⇒ Crucial for safety applications

► Usual estimation methods [Review in MI24a]

→ **Maximum likelihood**

→ **Cross-validation Mean Squared Error**

→ **Bayesian approaches**

➔ **New multi-objective estimation algorithm** for more reliable GP in [MI24b]

⇒ Reliable results observed in practice **up to a max dimension of $d = 20$**

Validation metrics → by cross-validation (LOO or K-fold CV)

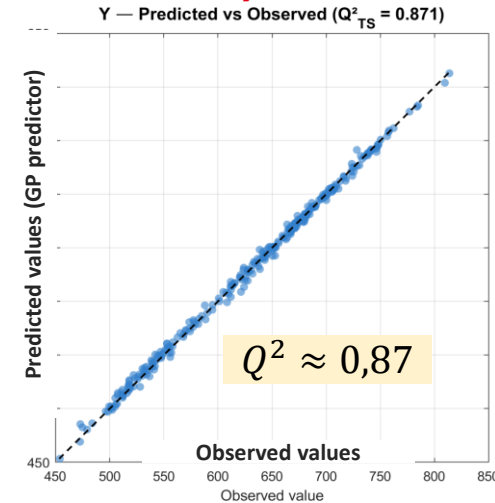
► Accuracy of the GPR predictor (only) [DIG⁺22]

$$Q^2 = 1 - \frac{MSE_{LOO}}{n^{-1} \sum_{i=1}^n (y_i - n^{-1} \sum_{i=1}^n y_i)^2}$$

$$\text{with } MSE_{LOO} = \frac{1}{n} \sum_{i=1}^n (y_i - \hat{y}_{-i}(\mathbf{x}^{(i)}))^2$$

where $\hat{y}_{-i}(\mathbf{x}^{(i)})$ is the GPR predictor in $\mathbf{x}^{(i)}$ when $(\mathbf{x}^{(i)}, y_i)$ is removed from the learning sample $\mathcal{D}_n$

→ The closer to 1 the Q^2 , the better the accuracy of the predictor.



► Accuracy of the whole GP predictive distribution [DIG⁺22, MI24a]

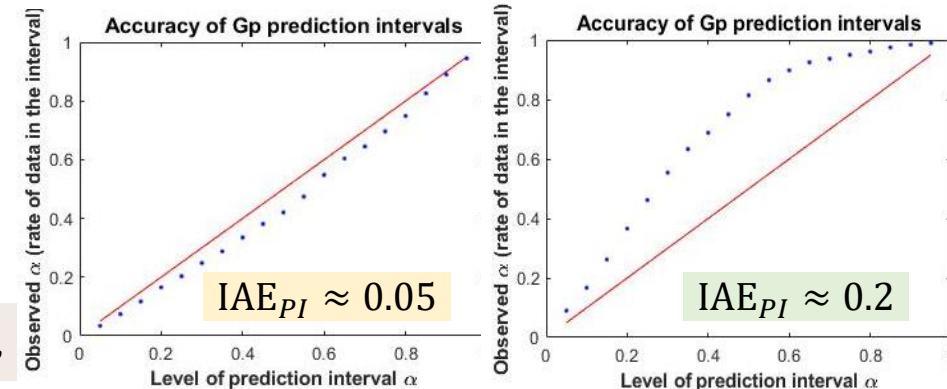
From empirical coverage function for $\alpha \in [0,1]$: $\hat{\Delta}(\alpha) = n^{-1} \sum_{i=1, \dots, n} \mathbf{1}\{y_i \in PI_{\alpha, -i}(\mathbf{x}^{(i)})\}$

with $PI_{\alpha, -i}(\mathbf{x}^{(i)})$ the α -level GP predictive interval for $\mathbf{x}^{(i)}$ when $(\mathbf{x}^{(i)}, y_i)$ is removed from $\mathcal{D}_n$

⇒ α -PI plot and integrated absolute error on $\hat{\Delta}(\alpha)$:

$$IAE_{PI} = \int_0^1 |\hat{\Delta}(\alpha) - \alpha|$$

→ The closer to 0 the IAE_{PI} , the better calibrated the predictive intervals.



Why and how to reduce dimension?



Objectives :

- ▶ Simplify and **improve reliability of GPR** hyperparameter estimation
- ▶ **Reduce computational burden** of Bayesian inference with GPR (UQ, calibration...)
- ▶ Find a **lower-dimensional input space** preserving **most output variability**, and potentially better suited to GPR stationarity assumptions

Under our applicative constraints

- ✓ Given data: single i.i.d. sample $\mathcal{D}_n$
- ✓ Small sample size: $n \approx 100$ to 1000
- ✓ Large input dimension: $d \approx 10$ to 100
- ✓ Gradient-free setting: black-box simulator

⇒ Proposed approach: HSIC screening + Active Subspaces (AS) + GPR [MI26]

How to build parsimonious GPR in large dimension?

1. Reminds on Gaussian Process Regression

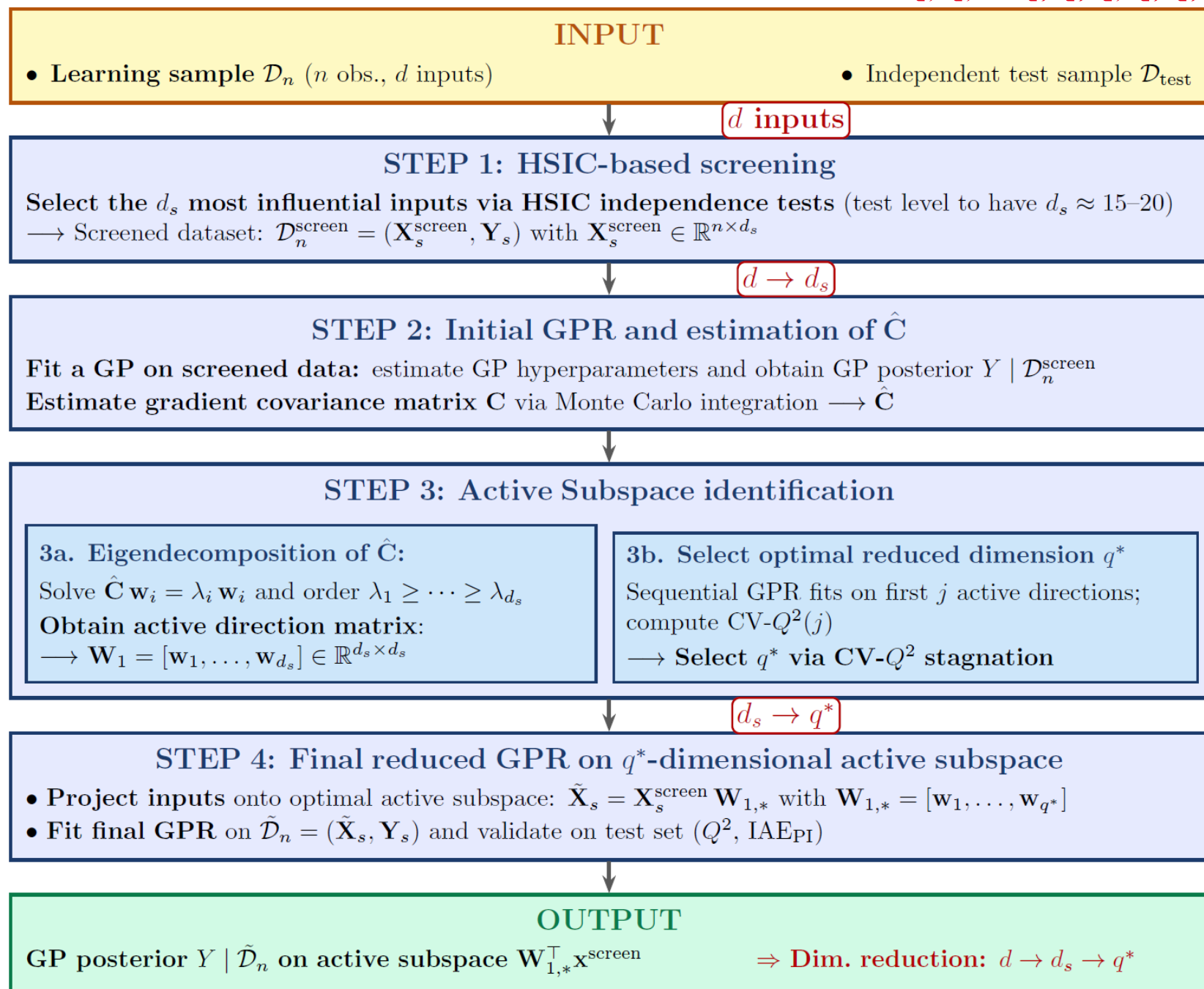
2. Proposed 4-step methodology → HSIC + AS + GPR

Proposed 4-step methodology

1st dim. reduction : $d \rightarrow d_s$

2nd dim. reduction : $d_s \rightarrow q^*$

Total dim. reduction: $d \rightarrow d_s \rightarrow q^*$



Step 1: HSIC-based screening

→ **1st dimension reduction** to obtain a reduced subset of $d_s \ll d$ inputs, to fit an initial GPR (step2)
→ Typically $d_s \lesssim 20$

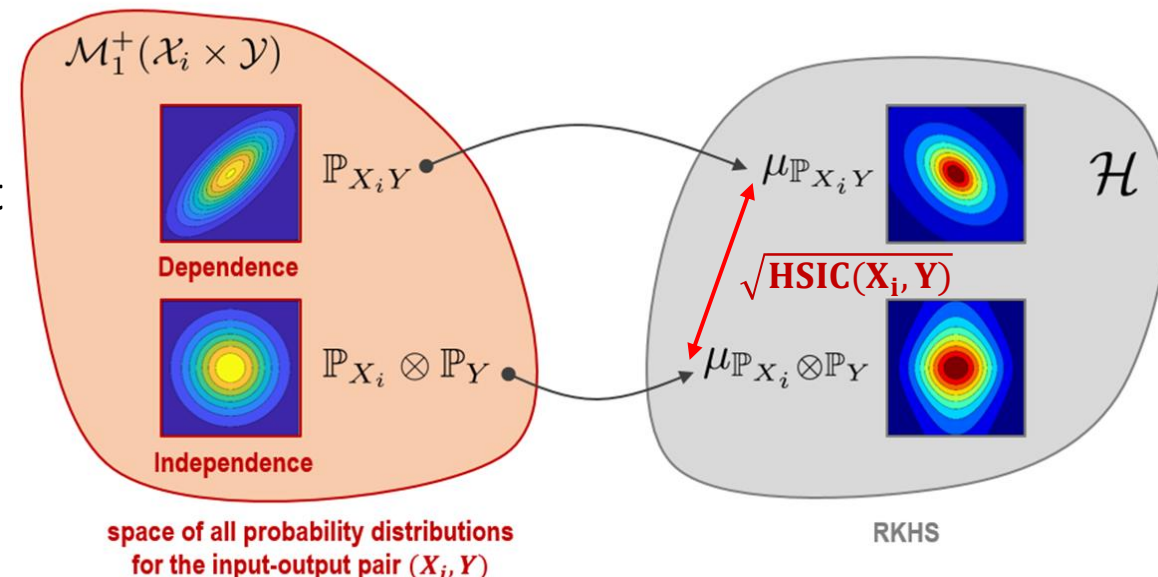
→ **HSIC-based SCREENING**: Surrogate-free + gradient-free + given-data and small-data context

► **HSIC*** [GFT+07, Dav15] measures dependence between X_i and Y via kernel embeddings in RKHS

- ✓ **Vanishes iff $X_i \perp\!\!\!\perp Y$** with characteristic kernels
→ **Independence test** for screening
- ✓ No distributional assumption, any input type
- ✓ Estimation from a **unique random sample**, efficient in practice from $n \sim 100$ (U-stat and V-stat estimators)

Standard kernel choice for real variables: **Gaussian kernels** with bandwidth = empirical standard deviation → characteristic kernels

$$HSIC(X_i, Y) = \left\| \mu_{\mathbb{P}_{X_i Y}} - \mu_{\mathbb{P}_{X_i} \otimes \mathbb{P}_Y} \right\|^2$$



***Hilbert-Schmidt Independence Criterion**

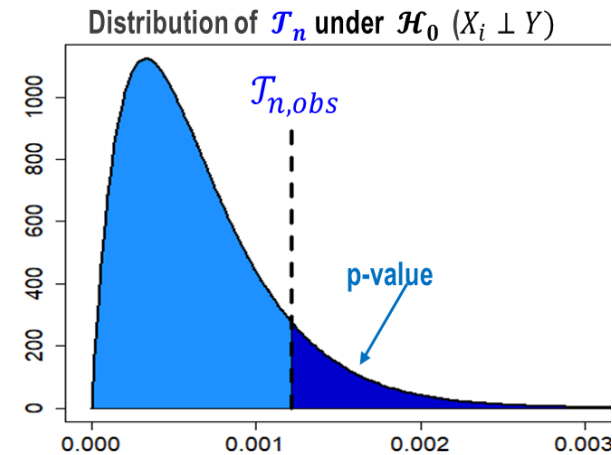
Extract from a presentation by G. Sarazin (CEA)

Step 1: HSIC-based screening

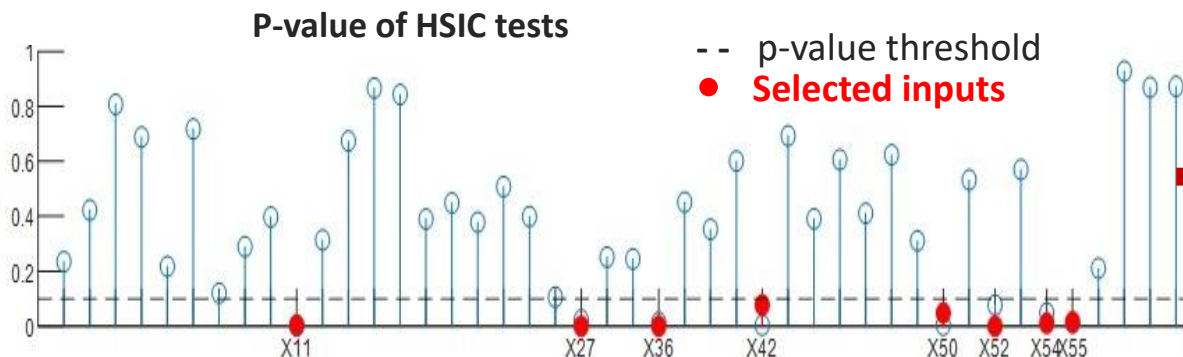
► **HSIC-based independence tests** [GFT+07]: $HSIC(X_i, Y) = 0 \Leftrightarrow X_i \perp\!\!\!\perp Y$ (with characteristic kernels!)

- ✓ Null hypothesis: $\mathcal{H}_0 : X_i \perp\!\!\!\perp Y$ against $\mathcal{H}_1 : X_i \not\perp\!\!\!\perp Y$
- ✓ Test statistics: $\mathcal{T}_n := n\widehat{HSIC}(X_i, Y)$
- ✓ Decision rule: $\mathcal{H}_0$ rejected iff $\mathcal{T}_n > q_{1-\alpha}$ with $q_{1-\alpha}$ the quantile of $\mathcal{T}_n$ under $\mathcal{H}_0$
 $\Leftrightarrow$ iff $p\text{-value} \leq \alpha$

→ Procedures to approximate the distribution of $\mathcal{T}_n$ under $\mathcal{H}_0$ [GFT+07, EM24]



⇒ Use for screening of inputs



Select inputs with $p\text{-value} \leq \alpha$ (e.g. $\alpha \approx 0.10\text{--}0.15$, or adaptive to retain $d_s \lesssim 20$)
→ d_s explicative inputs of GPR in Step 2

Step 2: build an initial GPR for AS estimation

- ▶ Fit a GPR on the reduced subset of d_s inputs (selected by HSIC-test)
 - ✓ **Covariance kernel choice for stable gradient estimation**: Matérn-5/2 or Gaussian kernel
 - The objective is to target average gradient structure, not local variation
 - $\mathcal{C}^1$ or $\mathcal{C}^2$ -diff kernel k required, depending on estimator of gradient cov. matrix $\hat{C}$ (→ next slide)
 - ✓ **Nugget regularization**: for training matrix only ($\tilde{K} = K + \tau^2 I_n$), NOT in base kernel k
 - Remove interpolator property (justified as $(d-d_s)$ non-selected inputs)
 - Smoother predictor: reduces sensitivity to local fluctuations
 - Preserves differentiability as it does not modify base kernel $k(\cdot, \cdot)$

Numerical illustration in [MI26]: **high-regularity kernels + nugget yield stable estimation of C**

Step 2: build an initial GPR for AS estimation

- Estimate the gradient covariance matrix $\mathbf{C}$ [Con15] defined for a model $\mathcal{M}$ as:

$$\mathbf{C} = \mathbb{E}_X[\nabla \mathcal{M}(X) \nabla \mathcal{M}(X)^\top] \quad \text{where } \nabla \mathcal{M}(\mathbf{x}) = \left(\frac{\partial \mathcal{M}}{\partial x_1}(\mathbf{x}), \dots, \frac{\partial \mathcal{M}}{\partial x_d}(\mathbf{x}) \right)^\top \in \mathbb{R}^d$$

Estimation from GPR: 2 closed-form estimators from posterior distribution of $\nabla Y | \mathcal{D}_n$ (which remains a GP)

- **Plug-in estimator** $\hat{\mathbf{C}}_{\text{plug-in}}$ (requiring $k \in \mathcal{C}^1$): only based on the posterior mean of $\nabla Y | \mathcal{D}_n$

$$\hat{\mathbf{C}}_{\text{plug-in}} = \mathbb{E}_X[\nabla \hat{y}(X) \nabla \hat{y}(X)^\top]$$

- **Bayesian estimator** $\hat{\mathbf{C}}_{\text{Bayes}}$ (requiring $k \in \mathcal{C}^2$): based on the full posterior distribution of $\nabla Y | \mathcal{D}_n$

$$\hat{\mathbf{C}}_{\text{Bayes}} = \mathbb{E}[\mathbb{E}_X[(\nabla Y | \mathcal{D}_n)(\nabla Y | \mathcal{D}_n)^\top]]$$

$$\rightarrow \hat{\mathbf{C}}_{\text{Bayes}} = \hat{\mathbf{C}}_{\text{plug-in}} + \mathbb{E}_X[\text{Cov}[\nabla Y | \mathcal{D}_n]]$$

Tests in [MI26]: negligible accuracy gain of $\hat{\mathbf{C}}_{\text{Bayes}}$, $\hat{\mathbf{C}}_{\text{plug-in}}$ is even more robust $\Rightarrow \hat{\mathbf{C}}_{\text{plug-in}}$ **default choice**

Step 3: identify Active Subspace

► Eigenvalue decomposition of $\hat{\mathbf{C}} \in \mathbb{R}^{d_s \times d_s}$: $\hat{\mathbf{C}} = \mathbf{W}\mathbf{\Lambda}\mathbf{W}^T$

with $\mathbf{W}$ the matrix of eigenvectors and $\mathbf{\Lambda}$ the diagonal matrix of ordered eigenvalues $\lambda_1 \geq \dots \geq \lambda_{d_s} \geq 0$

► Sequential GPR construction:

1. Project input data onto the j first active directions: $\mathbf{u}^{(j)} = \left(\mathbf{W}_1^{(j)}\right)^T \mathbf{x}$ with $\mathbf{W}_1^{(j)} \in \mathbb{R}^{d_s \times j}$
2. Fit a GPR surrogate on corresponding training dataset $\mathcal{D}_n^{(j)}$ (see details for kernel below*)
3. Estimate the predictivity coefficient $Q^2(j)$ via K-fold CV

► Selection of optimal dimension q^* of active subspace $\rightarrow$ Stagnation criterion of $Q^2(j)$

$\rightarrow$ Balance **parsimony against accuracy** of the final reduced GP

See [MI26]
for heuristics

*Kernel notes for sequential GPR:

- **Nugget** is part of the kernel $\tilde{k} = k + \tau^2 \delta$ (unlike Step 2) to absorb approximation error from restricting to $j < d_s$ directions
- **No constraint here on covariance regularity**, freely chosen to optimize predictive performance

Step 4: Final reduced GPR

- ▶ Final GPR construction on active subspace of dimension q^* (q^* first active directions)
- ▶ Validation by computation of Q^2 and IAE_{PI}

Large dataset → Test set available

- ✓ Compute directly Q^2 and IAE_{PI} on test set
→ Unbiased estimate of generalization performance

Small dataset → outer cross-validation $\cup$

- ✓ Fix d_s (Step 1) and q^* (Step 3) from full dataset
- ✓ Each fold: re-estimate Steps 2–3 on training portion, compute residuals on test portion
- ✓ Aggregate residuals → final CV Q^2 and IAE_{PI}
→ Slightly biased estimate of the error

Analytical benchmarks in [M126]: diverse variable subset embeddings, projections & intrinsic dimensionalities

⇒ Robustness of methodology confirmed

⇒ At equal reduced dim., AS+GPR identifies more performant subspaces than variable selection

How to build parsimonious GPR in large dimension?

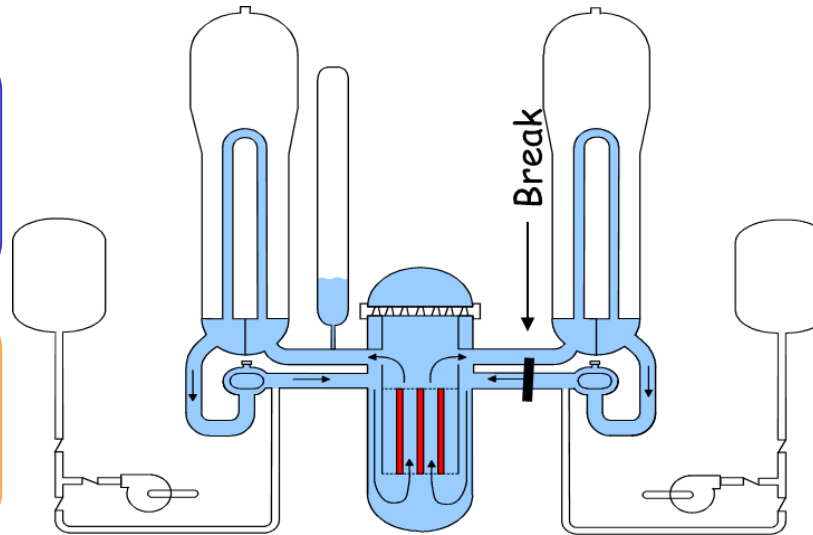
1. Reminds on Gaussian Process Regression
2. Proposed 4-step methodology → **HSIC + AS + GPR**
3. Illustration on thermal-hydraulic simulation testcase

Application: IB-LOCA Thermal-Hydraulic simulation



Pressurized Water Reactor scenario:
Loss of primary coolant accident due to a break in cold leg → **IB-LOCA accident**

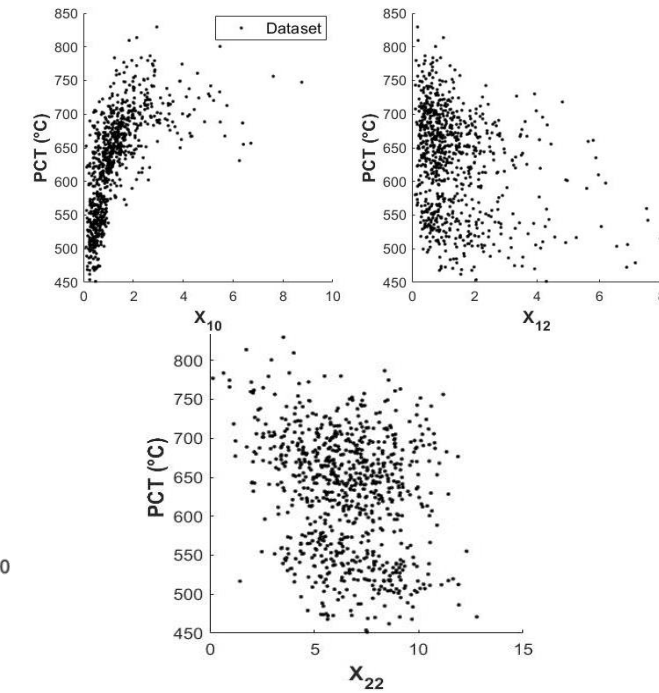
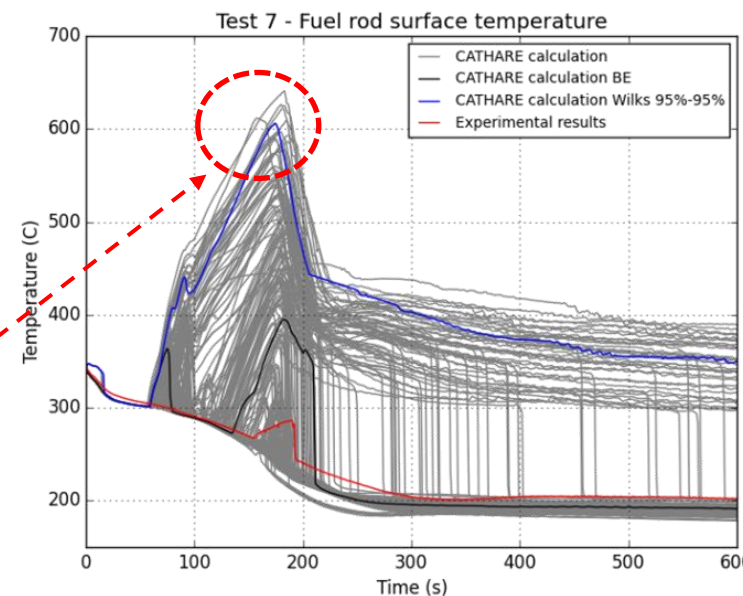
LSTF mock-up (1/48 scale) simulated with
thermal-hydraulic CATHARE2 code
→ CPU cost for one code run > 1 hour



Available dataset $\mathcal{D}_n$:
 $n = 1000$ Monte Carlo
simulations

$d = 27$ uncertain input variables X :
Critical flowrates, initial/boundary conditions, physical
equation coefficients, scenario parameters...

Output of interest
2nd peak of cladding temperature (PCT)
→ Critical safety indicator for licensing constraint



Application: IB-LOCA Thermal-Hydraulic simulation

3 approaches compared : constant mean + anisotropic stationary Matérn-5/2 for all GPRs

1. **Full-dim GPR** baseline: full $d=27$ inputs

2. **HSIC+GPR:**

✓ Standard HSIC screening with $\alpha = 0.1$ + GPR on selected inputs (kernel with nugget)

3. **AS+GPR [our 4-step methodology]** with:

✓ Step1: α adaptive $\rightarrow d_s=20$

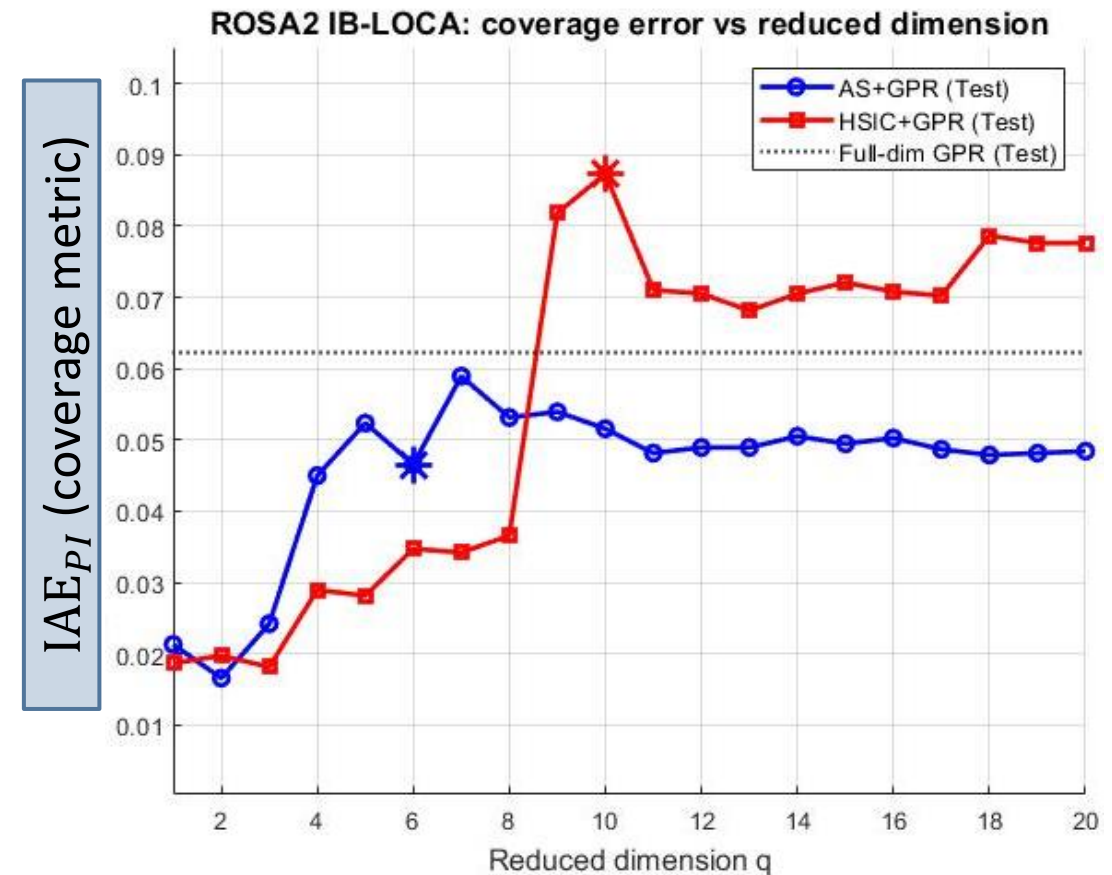
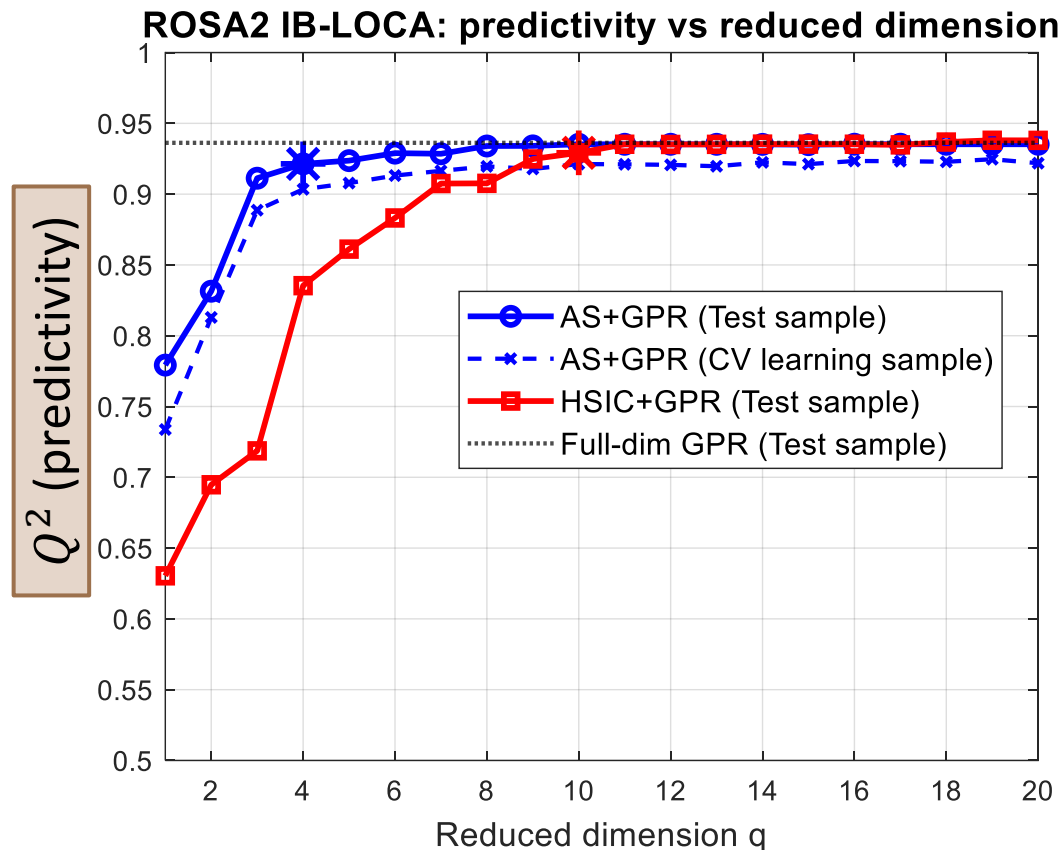
✓ Step2: nugget for training kernel matrix only, plug-in estimator $\hat{C}_{plug-in}$

✓ Step3: Q^2 -stagnation criterion of to select q^* (CV with 8 folds)

► Training set: $n_{train} = 700$ simulations

► Test set: $n_{test} = 300$

Application: IB-LOCA Thermal-Hydraulic simulation



Full-dim GPR : $\dim d = 27$

$Q^2 = 0.936$ $IAE_{PI} = 0.062$

→ High dimension, poor coverage

HSIC + GPR: $\dim d^* = 10$

$Q^2 = 0.929$ $IAE_{PI} = 0.087$

→ 63% dim. reduction, predictivity preserved but degraded coverage

AS + GPR: $\dim q^* = 6$

$Q^2 = 0.934$ $IAE_{PI} = 0.047$

→ 78% dim. reduction, no Q^2 loss and better coverage

Conclusions

► Methodology for an efficient and parsimonious GPR from Active Subspace:

- ✓ GPR benefits greatly from **preliminary HSIC-based screening**: it allows to fit a **1st GPR in medium dimension** (dim <20) to obtain a correct estimation of the **gradients** → **Reliable AS estimation**
- ✓ **Optimal AS reduction**: **heuristic approaches** proposed to select the optimal number of eigenvectors to keep, to balance GPR predictivity and dimension reduction
- ✓ **Final estimation of a low-dimensional GPR**: in reduced space derived from AS linear embedding
 - ⇒ Handles **high-dimensional, gradient-free** problems effectively
 - ⇒ More cost-effective use of **low-dim GPR for UQ** tasks in the Bayesian framework

► Extension to **modified active subspaces** [Lee19]

► Extension to **non-linear embedding** combined with selection or sparsity algorithm

[MI26] A. Marrel and B. Iooss (2026), Gaussian process and active subspace recipe for parsimonious metamodeling. Preprint (cea-05658731).

References 1/2

- [BW22] Binois, M., Wycoff, N., 2022. A survey on high-dimensional gaussian process modeling with application to Bayesian optimization. ACM Trans. Evol. Learn. Optim. 2.
- [Con15] Constantine, P. G. (2015). Active Subspaces. Philadelphia, PA: SIAM.
- [DL13] Damianou, A., Lawrence, N.D., 2013. Deep gaussian processes, in: Carvalho, C., Ravikumar, P. (Eds.), Proceedings of the Sixteenth International Workshop on Artificial Intelligence and Statistics, Scottsdale, AZ, USA. pp. 207–215.
- [Dav15] Da Veiga (2015). Global sensitivity analysis with dependence measures, Journal of Statistical Computation and Simulation, 85:1283-1305, 2015.
- [DIG⁺22] Demay, C., Iooss, B., Gratiet, L., and Marrel, A. (2022). Model selection for GP regression: an application with highlights on the model variance validation. QREI Journal, 38:1482-1500.
- [DRC13] Durrande, N., O., D.G., Roustant, Carraro, L., 2013. ANOVA kernels and RKHS of zero mean functions for model-based sensitivity analysis. Journal of Multivariate Analysis 155, 57–67.
- [DNR11] Duvenaud, D.K., Nickisch, H., Rasmussen, C.E., 2011. Additive gaussian processes, in: Advances in Neural Information Processing Systems, Curran Associates, Inc.. pp. 226–234.
- [EM24] El Amri and. Marrel (2024). More powerful HSIC based independence tests, extension to space filling designs and functional data. International Journal for Uncertainty Quantification 14(2): 69-98.
- [Gra21] B. Gramacy (2021) Gaussian Process Modeling, Design, and Optimization for the Applied Sciences. Chapman and Hall/CRC.
- [GFT⁺07] Gretton, A., Fukumizu, K., Teo, C., Song, L., Schölkopf, B. & Smola, A. (2007). A kernel statistical test of independence. Advances in Neural Information Processing Systems, 2007.
- [Lee19] Lee, M.R., 2019. Modified active subspaces using the average of gradients. SIAM/ASA Journal on Uncertainty Quantification 7, 53–66. doi:10.1137/17M1140662.

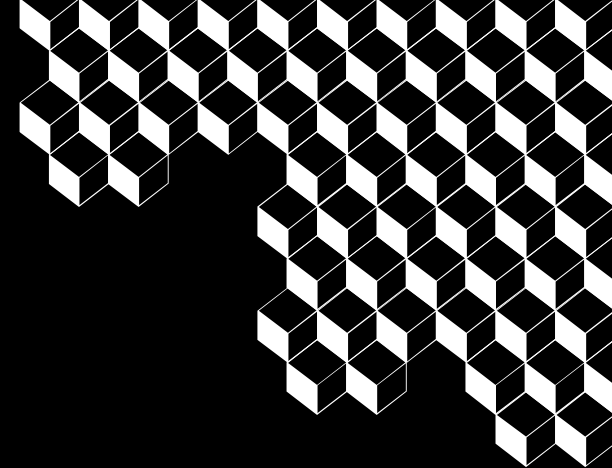
References 2/2



- [MI24a] A. Marrel and B. Iooss (2024), Probabilistic surrogate modeling by Gaussian process: A review on recent insights in estimation and validation, Reliability Engineering and System Safety, 247:110120.
- [MI24b] A. Marrel and B. Iooss (2024), Probabilistic surrogate modeling by Gaussian process: A new estimation algorithm for more robust prediction, Reliability Engineering and System Safety, 247:110094.
- [RW06] C.E. Rasmussen and C.K.I. Williams (2006). Gaussian processes for machine learning. MIT Press.
- [RWB+20] Raponi, E., Wang, H., Bujny, M., Boria, S., Doerr, C., 2020. High dimensional bayesian optimization assisted by principal component analysis, in: Parallel Problem Solving from Nature – PPSN XVI, Springer. pp. 169–183.
- [SMD+23] Sarazin, G., Marrel, A., Da Veiga, S. & Chabridon (2023). New insights into the feature maps of Sobolev kernels: application in global sensitivity analysis. <https://cea.hal.science/cea-04320711>.
- [TBG16] Tripathy, R., Bilonis, I., Gonzalez, M., 2016. Gaussian processes with built-in dimensionality reduction: Applications to high-dimensional uncertainty propagation. Journal of Computational Physics 32, 191–223.
- [TB19] Tripathy, R., Bilonis, I., 2019. Deep active subspaces: A scalable method for high-dimensional uncertainty propagation, in: ASME 2019 International Design Engineering Technical Conferences and Computers and Information in Engineering Conference, American Society of Mechanical Engineers. p. V001T02A074.
- [WBW21] N. Wycoff, M. Binois, S. M. Wild (2021). Sequential Learning of Active Subspaces. Journal of Computational and Graphical Statistics, 30(4), 1224–1237.
- [WHS+17] Wilson, A.G., Hu, Z., Salakhutdinov, R., Xing, E.P., 2017. Deep kernel learning. Journal of Machine Learning Research 18, 1–33. Available at: <https://arxiv.org/pdf/1511.02222>
- [WHZ+16] Wang, Z., Hutter, F., Zoghi, M., Matheson, D., De Freitas, N., 2016. Bayesian optimization in a billion dimensions via random embeddings. Journal of Artificial Intelligence Research 55, 361–387.



Appendix



Why choosing Active Subspaces (AS)?

Approach	Advantages ✓	Limitations ✗
Linear embeddings <i>AS, PCA, random projections, linear manifolds [Con15, TBG16, WHZ+16, RWB+20]</i>	Simple, interpretable, small data-friendly, theoretical guarantees	May degrade with strong nonlinearities
Non-linear embeddings coupled with GPR <i>Deep kernel learning, deep GP, VAE... [WHS+17, DL13, TB19]</i>	Expressive for complex functions	Large data needed Expensive Low interpretability
Additive / ANOVA-type GP <i>Additive GP, sparse ANOVA kernels... [DRC13, DNR11] + Review of [BW22]</i>	Effective for sparse additive structures	Complex model selection Not general-purpose

Adapted to our applicative settings and objectives

Our contribution [MI26]: a **practical 4-step methodology combining GPR with AS in gradient-free settings**, robustly estimating AS and selecting the optimal number of active directions

Estimation of GP hyperparameters (θ, τ)

⇒ How to robustly estimate the **hyperparameters** $(\theta, \tau) \in \mathbb{R}^{+,d+1}$ from the learning sample?

➔ For good **predictivity** + **reliable prediction intervals** ⇒ Crucial for safety applications

► Usual estimation methods [Review in MI24a]

→ **Maximum likelihood**

→ **Cross-validation Mean Squared Error**

→ **Bayesian approaches**

Ill-posedness, problem of **flatness** of
functions to be minimized

CPU ++, **delicate choice of priors**
Except RobustGAsp method

➔ **New multi-objective estimation algorithm** for more reliable GP in [MI24b]

⇒ Reliable results observed in practice **up to a max dimension of $d = 20$**

Step 1: HSIC-based screening

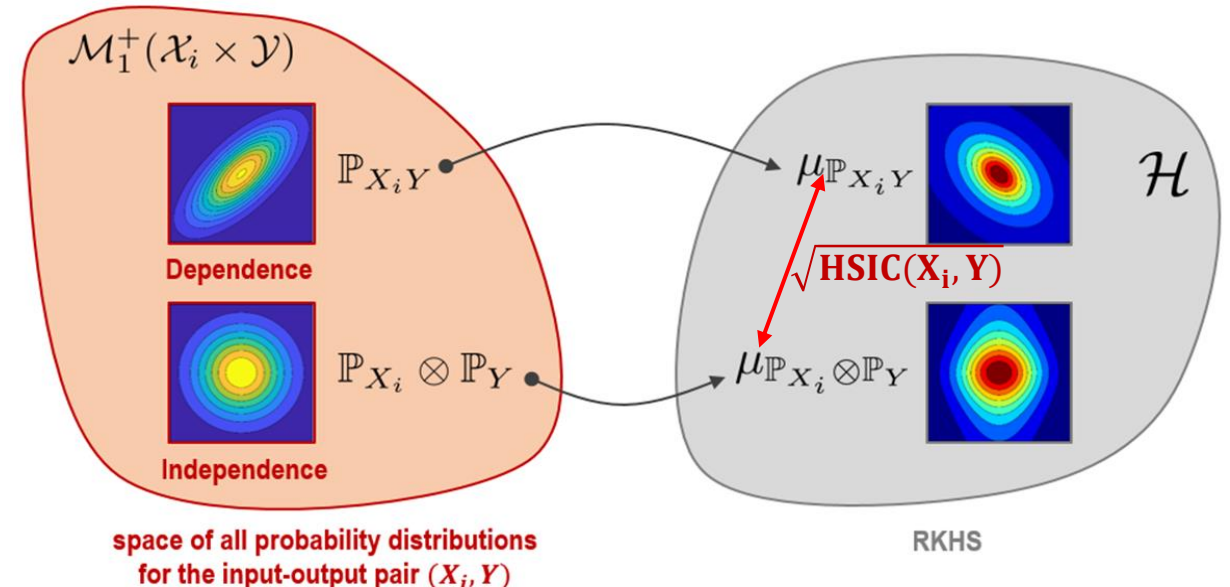
➔ **1st dimension reduction** to obtain a reduced subset of $d_s \ll d$ inputs, enabling reliable initial GPR estimation in Step 2 → **Typically $d_s \lesssim 20$**

➔ **HSIC-based SCREENING:** Surrogate-free + Gradient-free + Adapted to given-data and small-data context

► **HSIC*** ([GFT+07]) measure dependence between X_i and Y via kernel embeddings in RKHS

$$HSIC(X_i, Y) = MMD^2(P_{X_i Y}, P_{X_i} \otimes P_Y) = \left\| \mu_{P_{X_i Y}} - \mu_{P_{X_i} \otimes P_Y} \right\|^2$$

- ✓ **HSIC**: Estimation of from a **unique random sample**, efficient in practice from $n \sim 100$
- ✓ HSIC "summarizes" the cross-cov between feature maps $\Rightarrow$ Large panel of input-output dependency captured [Dav15]



Dealing with the large input dimension

► HSIC-based ranking [Dav15] :

$$R^2_{HSIC} = \frac{HSIC(X,Y)}{\sqrt{HSIC(X,X)HSIC(Y,Y)}} \Rightarrow R^2_{HSIC} \in [0,1] \text{ for easier interpretation}$$

$$\text{Influence}(X_{[1]}) > \text{Influence}(X_{[2]}) > \dots > \text{Influence}(X_{[d]})$$

Where order $[\cdot]$ is such that $\hat{R}^2_{H,X_{[1]}} > \hat{R}^2_{H,X_{[2]}} > \dots > \hat{R}^2_{H,X_{[d]}}$

⇒ Use for ranking of inputs

Inputs ordered by degree of influence



Can be used for **more robust sequential GPR estimation**

⇒ “forward” estimation of GPR hyperparameters: successive inclusion of ordered inputs

HSIC review: a kernel-based GSA method

► **MMD²** applied between $P_{X_i Y}$ and $P_{X_i} \otimes P_Y \Rightarrow HSIC(X_i, Y)_{\mathcal{H}_{X_i}, \mathcal{H}_Y}$

$\mathcal{H}_{X_i}$ and $\mathcal{H}_Y$ **RKHS** associated to X_i and Y , resp :

Kernel $k_{X_i}: \mathcal{X}_i \times \mathcal{X}_i \rightarrow \mathbb{R}$ with feature space $\mathcal{H}_{X_i}$ and feature map φ_{X_i}

Kernel $k_Y: \mathcal{Y} \times \mathcal{Y} \rightarrow \mathbb{R}$ with feature space $\mathcal{H}_Y$ and feature map φ_Y

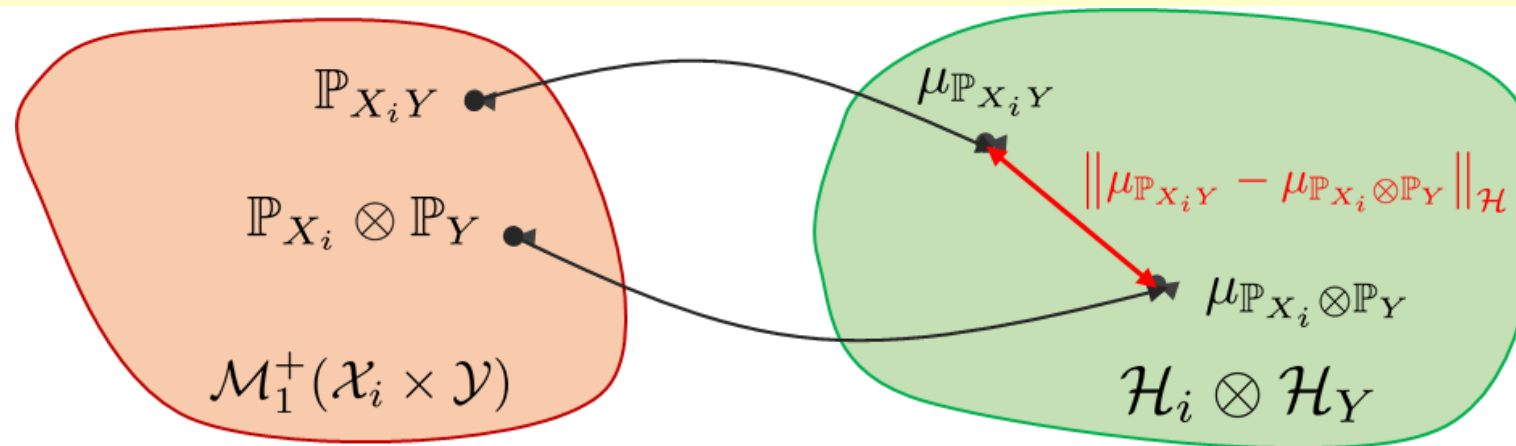
$$K_{X_i}(x, x') = \langle \varphi_{X_i}(x), \varphi_{X_i}(x') \rangle_{\mathcal{H}_{X_i}} \text{ and } K_Y(y, y') = \langle \varphi_Y(y), \varphi_Y(y') \rangle_{\mathcal{H}_Y}$$

kernel defines the inner product in the RKHS

HSIC = distance in the RKHS between the images of the two distributions of interest

Gretton et al. [2005]

$$\Rightarrow HSIC(X_i, Y)_{\mathcal{H}_{X_i}, \mathcal{H}_Y} = MMD_{\mathcal{H}_{X_i}, \mathcal{H}_Y}^2(P_{X_i Y}, P_{X_i} \otimes P_Y) = \left\| \mu_{\mathbb{P}_{X_i Y}} - \mu_{\mathbb{P}_{X_i} \otimes \mathbb{P}_Y} \right\|_{\mathcal{H}_{X_i}, \mathcal{H}_Y}^2$$



Space of all probability distributions
for the input-output pair

Tensorized RKHS

Extracted from G. Sarazin's (CEA) slides

HSIC review: a kernel-based GSA method

► **MMD²** applied between $P_{X_i Y}$ and $P_{X_i} \otimes P_Y \Rightarrow HSIC(X_i, Y)_{\mathcal{H}_{X_i}, \mathcal{H}_Y}$

$\mathcal{H}_{X_i}$ and $\mathcal{H}_Y$ **RKHS** associated to X_i and Y , resp :

Kernel $k_{X_i}: \mathcal{X}_i \times \mathcal{X}_i \rightarrow \mathbb{R}$ with feature space $\mathcal{H}_{X_i}$ and feature map φ_{X_i}

Kernel $k_Y: \mathcal{Y} \times \mathcal{Y} \rightarrow \mathbb{R}$ with feature space $\mathcal{H}_Y$ and feature map φ_Y

$$K_{X_i}(x, x') = \langle \varphi_{X_i}(x), \varphi_{X_i}(x') \rangle_{\mathcal{H}_{X_i}} \text{ and } K_Y(y, y') = \langle \varphi_Y(y), \varphi_Y(y') \rangle_{\mathcal{H}_Y} \quad \text{kernel defines the inner product in the RKHS}$$

HSIC = distance in the RKHS between the images of the two distributions of interest

Gretton et al. [2005]

$$\Rightarrow HSIC(X_i, Y)_{\mathcal{H}_{X_i}, \mathcal{H}_Y} = MMD_{\mathcal{H}_{X_i}, \mathcal{H}_Y}^2(P_{X_i Y}, P_{X_i} \otimes P_Y) = \|C_{X,Y}\|_{HS}^2$$

With $C_{X,Y}$ the **covariance operator** between features maps:

$$C_{X,Y} = \mathbb{E}_{X,Y} [\varphi_X(X) \otimes \varphi_Y(Y)] - \mathbb{E}_X [\varphi_X(X)] \otimes \mathbb{E}_Y [\varphi_Y(Y)]$$

HSIC "summarizes" the cross-cov between feature maps

$\Rightarrow$ Large panel of input-output dependency can be captured.

HSIC review: a kernel-based GSA method

► Characteristic kernels and RKHS $\Rightarrow$ Injective canonical feature map

$\Rightarrow$ Equivalence to independence: $HSIC(X, Y) = 0 \Leftrightarrow X \perp Y$

Ex: Gaussian Kernel

$$k(x_i, x'_i) = \exp\left(-\frac{(x_i - x'_i)^2}{2\lambda^2}\right)$$

► Estimation: Kernel Trick $\Rightarrow$ Feature map linked to kernel function

Very simple M-C estimator from a n -sample of simulations $(X_i^{(j)}, Y^{(j)})_{1 \leq j \leq n}$

$$\widehat{HSIC}(X_i, Y) = \frac{1}{n-1} \text{Tr}(K_i H L H)$$

$$\text{where } H = I_n - \frac{1}{n}, K_i = \left(k_i \left(X_i^{(j)}, X_i^{(j')} \right) \right)_{1 \leq j, j' \leq n} \text{ and } L = \left(k \left(Y^{(j)}, Y^{(j')} \right) \right)_{1 \leq j, j' \leq n}$$

► Statistical Properties:

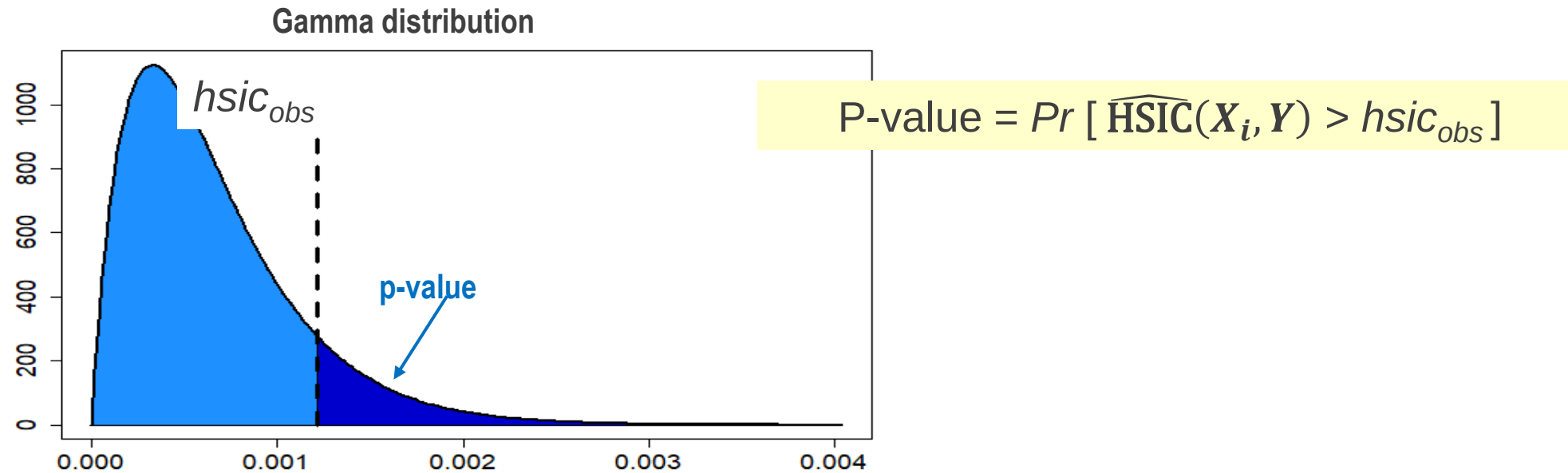
- Asymptotically unbiased, variance of order $O(1/n)$
- If $X \perp Y$, $n\widehat{HSIC}(X, Y)$ converges asymptotically to a Gamma distribution

HSIC review: a kernel-based GSA method

HSIC-based independence tests for screening

How to have the distribution $n\widehat{\text{HSIC}}(X_i, Y)$ under $\mathcal{H}_0$ to compute p -value?

- If n large: asymptotic test based on approximation with Gamma law (Gretton et al. (2008))
- If n small: Permutation-based approximation (De Lozzo & Marrel [2016a], Meynaoui [2019], El Amri & Marrel [2021a])



Interpretation of p -value for a level α ($\alpha = 5\%$ or 10%) for screening:

- $p\text{val} < \alpha$ $\Rightarrow H_0$ (Independence) rejected $\Rightarrow X_i$ is significantly influential

GPR validation

► Validation criteria computed by cross-validation (LOO or K-fold CV) [DIG⁺21]

→ Accuracy of the whole GP conditional distribution [DIG⁺21,ABG23]

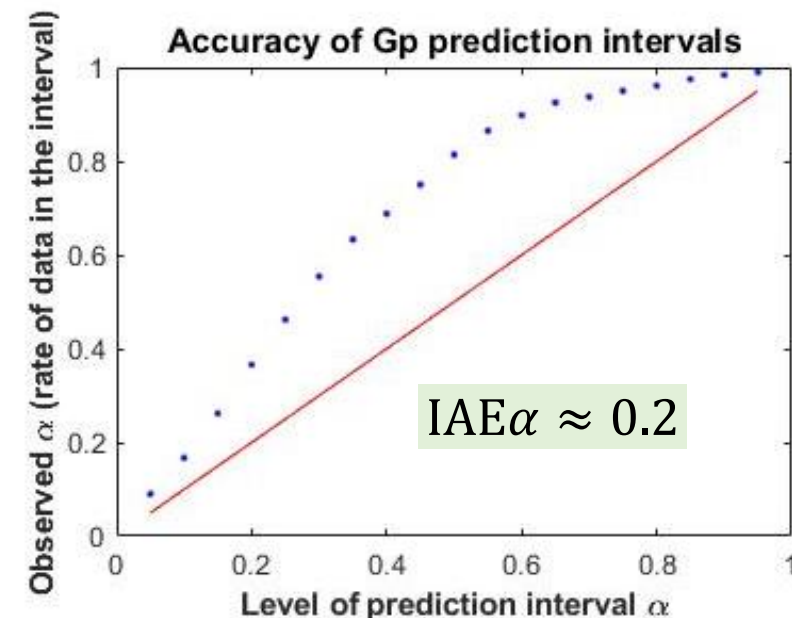
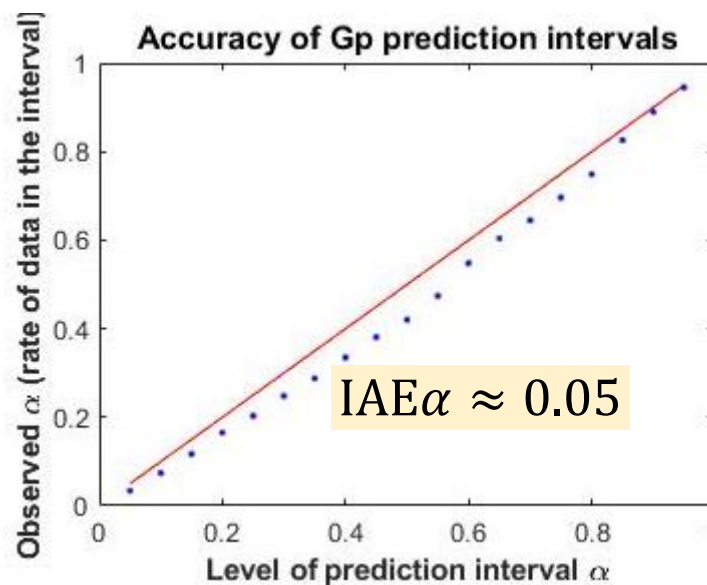
From empirical coverage function for $\alpha \in [0,1]$:
$$\hat{\Delta}(\alpha) = \frac{1}{n} \sum_{i=1}^n \mathbf{1}\{y_i \in PI_{\alpha,-i}(\mathbf{x}^{(i)})\}$$

with $PI_{\alpha,-i}(\mathbf{x}^{(i)})$ the α -level GP prediction interval for $\mathbf{x}^{(i)}$ when $(\mathbf{x}^{(i)}, y_i)$ is removed from learning sample

⇒ α -PI Plot

⇒ Integrated Absolute Error on $\hat{\Delta}(\alpha)$
[MI24a]

$$IAE\alpha = \int_0^1 |\hat{\Delta}(\alpha) - \alpha|$$

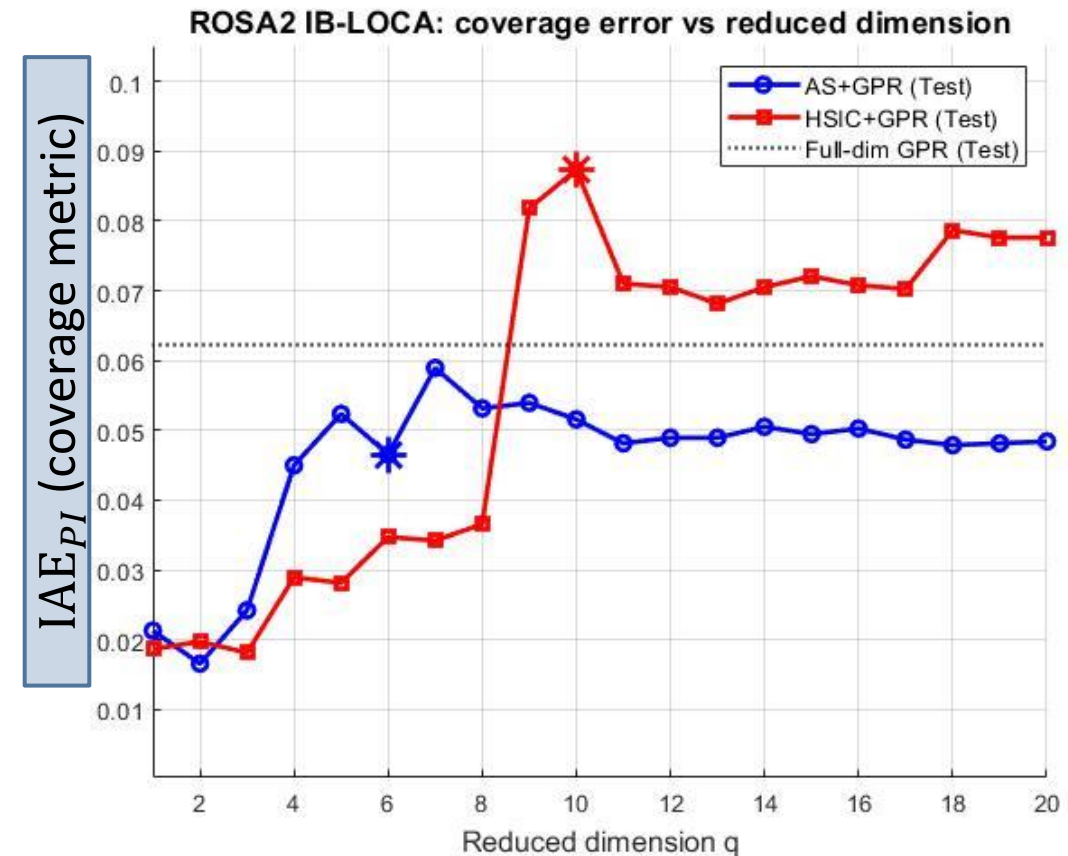
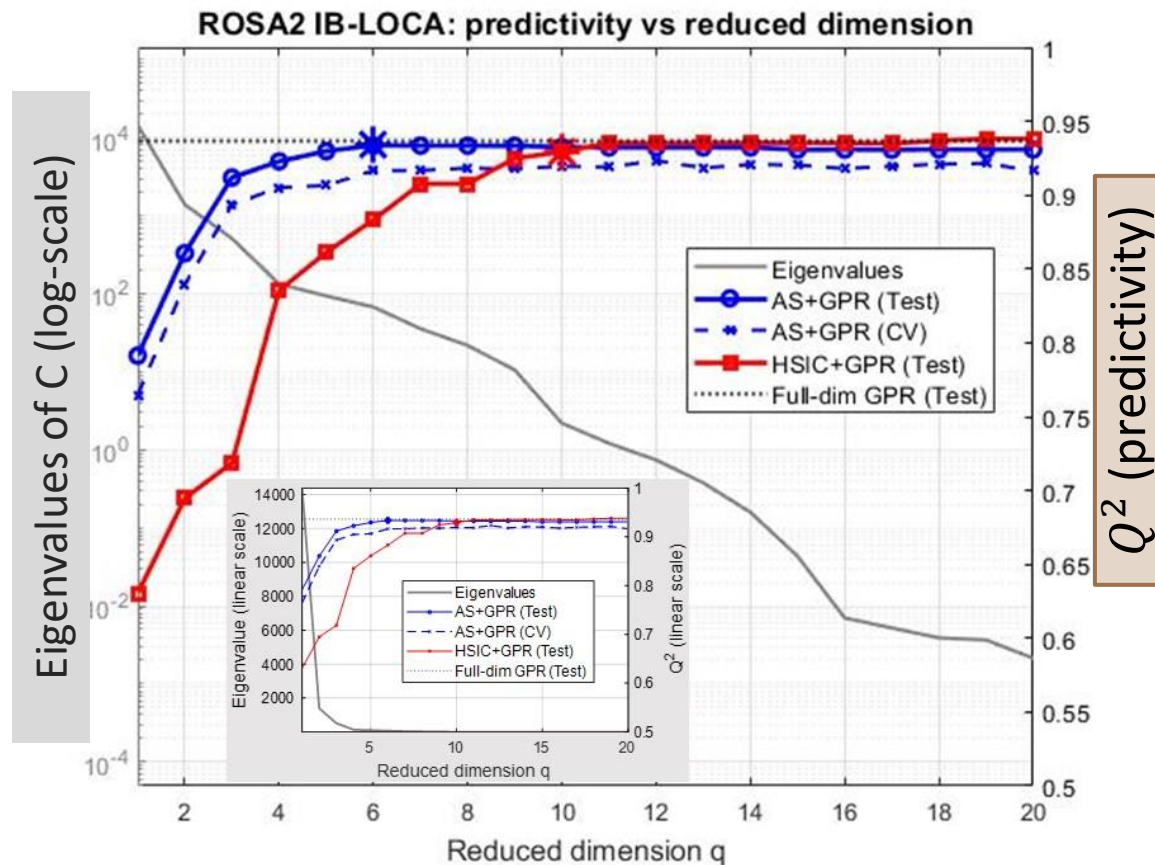


GPR validation

► One message: Joint interpretation of Q^2 and $IAE\alpha$! [MI24a]

Criterion	Values	Interpretation
$IAE\alpha$	Value close to 0	<u>Only if Q^2 is also high</u> , reliable predicted confidence intervals
	High values, close to 0.5 e.g.	Unreliable predicted prediction intervals ("underconfident" or "overconfident" model) ⇒ Explanation from cross-interpretation with Q^2 and α-PI plot

Application: IB-LOCA Thermal-Hydraulic simulation



Full-dim GPR : $\dim d = 27$

$Q^2 = 0.936$ $IAE_{PI} = 0.062$

→ High dimension, poor coverage

HSIC + GPR: $\dim d^* = 10$

$Q^2 = 0.929$ $IAE_{PI} = 0.087$

→ 63% dim. reduction, predictivity preserved but degraded coverage

AS + GPR: $\dim q^* = 6$

$Q^2 = 0.934$ $IAE_{PI} = 0.047$

→ 78% dim. reduction, no Q^2 loss and better coverage

Illustration of methodology for parsimonious GPR

Becker Model $\mathcal{M}_{\text{Becker}}$ [Bec20] with order 2 interaction effects:

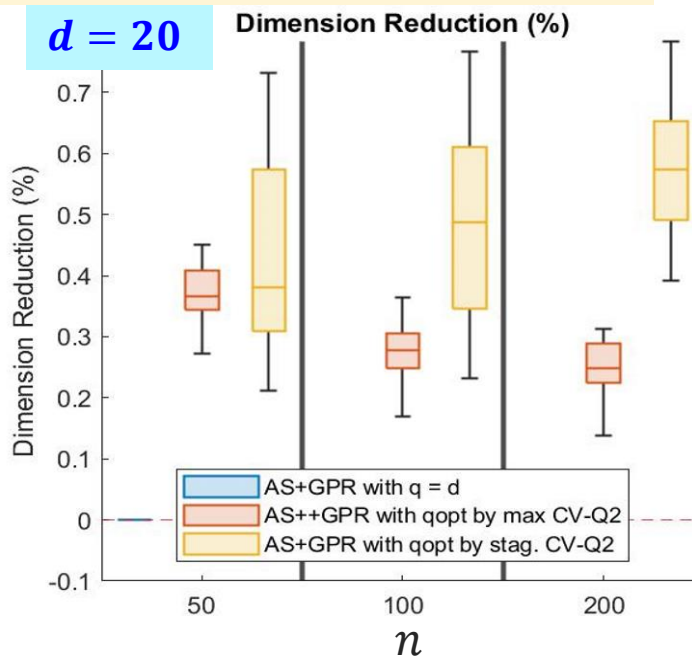
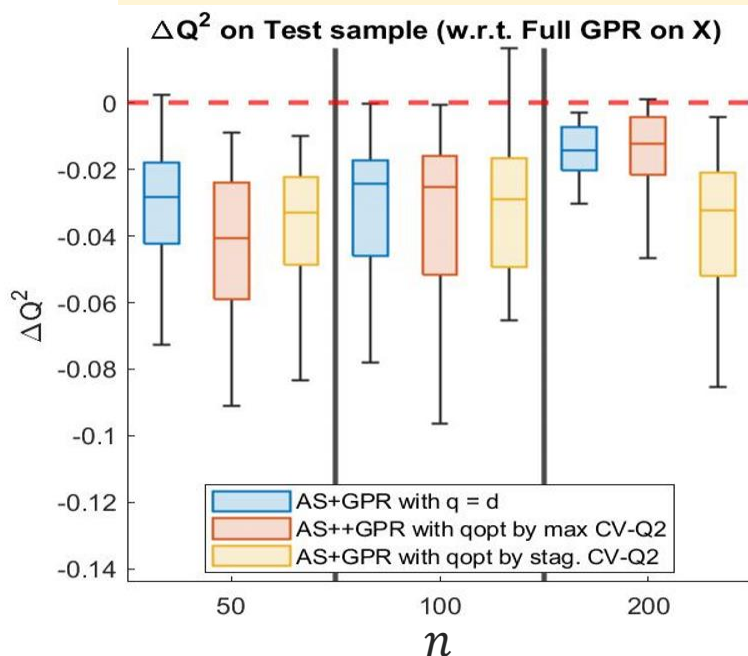
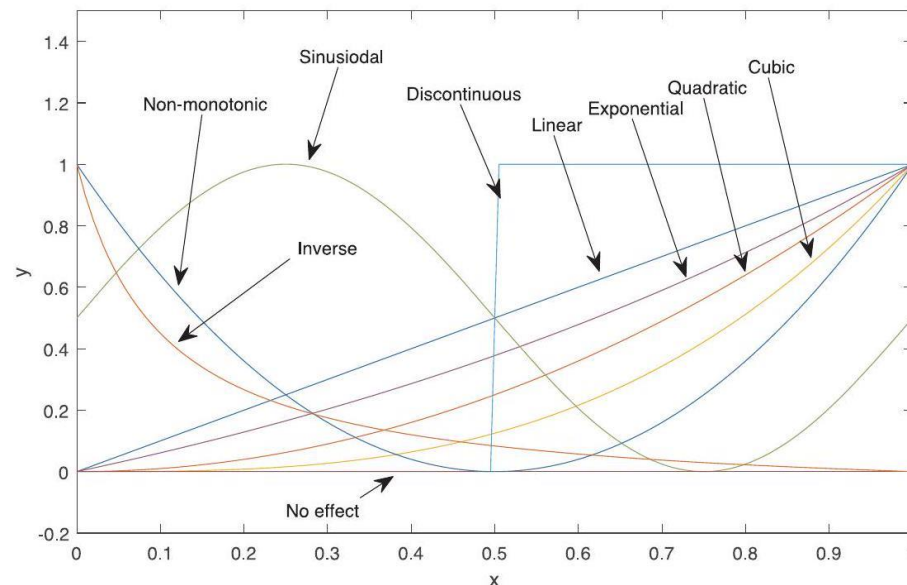
$$\mathcal{M}_{\text{Becker}}(\mathbf{x}, \mathbf{u}, \mathbf{V}, \Theta) = \sum_{i=1}^d a_i f^{u_i}(x_i) + \sum_{i=1}^{d_2} b_i f^{u_{V_{i,1}}}(x_{V_{i,1}}) f^{u_{V_{i,2}}}(x_{V_{i,2}}) \text{ with } \Theta = \{a, b\}$$

→ Built-upon a collection of 9 elementary functions f^k

→ Random sampling of parameters $\{u, V, \Theta\} \rightarrow$ Randomized function $\mathcal{M}_{\text{Becker}}$

$d = 20$ independent uniform inputs: $\forall i \in \llbracket 1, 20 \rrbracket, X_i \sim \mathcal{U}([0, 1])$

For each size n , 20 random draws of $\mathcal{M}_{\text{Becker}}$ functions.
(On average 70% of primary effects + 30% of interaction effects)



Results for GPR with Matern 5/2 covariance

- **Predictivity Preserved** : With q selected by maximizing or stagnating Q^2 (estimated by CV*)
→ reduced GPR explains **less than 5% less variance** compared to the high-dimensional GPR
- **Significant Dimension Reduction**: On average, 25% (respectively 50%) with heuristic based on maximization (respectively stagnation) of Q^2 .

CV* : cross-validation