

Validation of Machine Learning techniques in joint default assessment

Patrizia Semeraro¹

joint work with

Elisa Luciano²

Edoardo Fadda³

Bayes Business School

London, November 22, 2022.

¹Department of Mathematical Sciences G. Lagrange, Politecnico di Torino

²ESOMAS Department and Collegio Carlo Alberto, Università di Torino

³Department of Mathematical Sciences G. Lagrange, Politecnico di Torino

Motivation and aim

Motivation

ML techniques have recently been considered as alternatives to traditional LR models to predict default probabilities in the literature.

Aim

Empirically measure the model risk associated to the choice of a ML method to predict individual defaults.

Three levels:

- ① We compare different choices of a ML technique to estimate individual default probabilities of a portfolio of obligors;
- ② we study its impact on exchangeable portfolios:
 - we study the impact of estimating p with different ML methods on the risk of an exchangeable portfolio when the dependence structure is unspecified;
 - we study its impact on the VaR of the distribution of one exchangeable credit portfolio.

Framework

This presentation is partially based on the previous works:

Theoretical bounds for risk measures

Fontana, Roberto, Elisa Luciano, and Patrizia Semeraro. "Model risk in credit risk." *Mathematical Finance* 31.1 (2021): 176-202.

Preliminary analysis on ML performance on credit cards data

M. Doria, E. Luciano and P. Semeraro. "Machine Learning techniques in joint default assessment," <https://arxiv.org/abs/2205.01524>, 3 May 2022.

The model

Dependent defaults model

- $\mathbf{Y} = (Y_1, \dots, Y_d)$ random vector of default indicators, referring to d obligors over a fix time horizon T .
- The **loss** of portfolio P is given by

$$L = \sum_{i=1}^d w_i Y_i, \quad w_i = \frac{1}{d} \quad i = 1, \dots, d,$$

where $w_i \in (0, 1]$, $\sum_{i=1}^d w_i = 1$.

- We assume **exchangeability** among defaults and we focus on the case of homogeneous exposures $w_i = \frac{1}{d}$, $i = 1, \dots, d$, so that $L = \frac{S}{d}$. The **number of defaults** is given by

$$S = \sum_{i=1}^d Y_i$$

We have

$$\mathbf{Y} \leftrightarrow S \tag{1}$$

The classes $\mathcal{E}_d(p)$, and $\mathcal{S}_d(p)$

We call $\mathcal{S}(p)$ the classes of distributions of the number of defaults for exchangeable vectors in $\mathcal{E}(p)$.

Let $\mathbf{X} \in \mathcal{E}_d(p)$, i.e $\mathbf{Y}$ is a vector of Bernoulli $B(p)$ exchangeable variables, $f_p(\mathbf{y}) = f_p(\sigma(\mathbf{y}))$ for any $\sigma \in \mathcal{P}_d$

Number of defaults

Let now $\mathcal{S}_d(p)$ be the class of distributions of S_d , where

$$S := S_d = \sum_{i=1}^d X_i,$$

and $\mathbf{p}_s = (p_0, \dots, p_d)$, with $p_j = P(S_d = j)$.

$$\mathcal{E}(p) \leftrightarrow \mathcal{S}_d(p) \tag{2}$$

The class $\mathcal{S}_d(p)$

Theorem

The following holds. $S_d \in \mathcal{S}_d(p)$ iff there exist $\lambda_1, \dots, \lambda_{n_p} \geq 0$ summing up to 1 such that

$$\mathbf{p}_S = \sum_i^{n_p} \lambda_i \mathbf{r}_S^i, \quad (3)$$

where $\mathbf{r}_S^i$ are the extremal densities and n_p is the number of extremal densities.

Extremal distributions

The extremal distributions of $\mathcal{S}_d(p)$ have support on at most two points and we can find them analytically.

The parametrical model

We represent $\mathbf{Y}$ and S through an **exchangeable Bernoulli mixture model**.

Definition

Given a random variable Q , the random vector $\mathbf{Y} = (Y_1, \dots, Y_d)^T$ follows an exchangeable Bernoulli mixture model with mixing variable Q with support on $[0, 1]$, if conditional on Q the default indicator $\mathbf{Y}$ is a vector of independent Bernoulli random variables with $\mathbb{P}(Y_i = 1|Q) = Q$.

We assume that the mixing variable Q follows a beta distribution with parameters a and b , i.e. $Q \sim \beta(a, b)$: the beta mixing-model.

Model features

Exchangeable Bernoulli mixture model

The unconditional marginal default probability becomes

$$\mathbb{P}(Y_i = 1) = \int_0^1 q dG(q), \quad (4)$$

where $Q \sim G(q)$, the unconditional probability mass function (pmf) $p_Y(\mathbf{y})$ of $\mathbf{Y}$ becomes:

$$p_Y(\mathbf{y}) = \mathbb{P}(\mathbf{Y} = \mathbf{y}) = \int_0^1 q^{\sum_{i=1}^d y_i} (1 - q)^{d - \sum_{i=1}^d y_i} dG(q). \quad (5)$$

Finally, the distribution $p_S(k)$ of the number of defaults S becomes:

$$p_S(k) = \binom{d}{k} \int_0^1 q^k (1 - q)^{d-k} dG(q). \quad (6)$$

We consider $Q_h \sim \beta(a, b)$, so that S follows a $\beta\text{Bin}(d, a, b)$ distribution.

Model features

Exchangeable Bernoulli mixture model

- Cross moments of $\mathbf{Y}$:

$$\pi_k = E[Y_{i_1} \cdots Y_{i_k}], \quad \{i_1, \dots, i_k\} \subset \{1, \dots, d\} \quad 1 \leq k \leq d$$

Relevant quantities:

- $\pi_1 = E[Y_i] = P(Y_i = 1) := p$ marginal default probability
- $\rho = \rho(Y_i, Y_j) = \frac{\pi_2 - p^2}{p(1-p)} \quad i \neq j$
- The cross moments of $\mathbf{Y}$ are the moments of the mixing variable Q_h , i.e.:

$$\pi_k = \mathbb{E}[Q_h^k]$$

in particular $\pi_1 = \mathbb{E}[Q_h] = p$.

- S rv modelling the number of defaults, with distribution

$$P(S = k) = \sum_{i=0}^{d-k} (-1)^i \frac{d!}{i!k!(d-k-i)!} \pi_{k+i} \quad (7)$$

Portfolio risk

Risk measures

Let S be a random variable representing the number of defaults with finite mean. The **VaR $_{\alpha}$** at level α is defined by

$$\text{VaR}_{\alpha}(S) = \inf\{y \in \mathbb{R} : P(S \leq y) \geq \alpha\}$$

Default probability p and default correlation ρ are the main factors influencing the tails of S .

Aim

Model risk of exchangeable portfolios: effect of choosing a ML method to estimate p on VaR -bounds.

Model risk associated to a single model: effect of choosing a ML method to estimate the β Binomial distribution of defaults on the VaR.

Effect of p on VaR and ES

Effect of p on S

We have analytical sharp bounds for the VaR and ES in the classes $\mathcal{E}_d(p)$.

- Luciano Fontana and Semeraro (2020) proved That The sharp bounds for VaR are on the extremal points.
- ES is a convex riks of measure, its bounds are on the minimal convex sums extremal points (unique in the exchangeable case) and on the upper Fréchet bound.

Extremal generators of $\mathbf{f}_p \in \mathcal{E}_d(p)$ and $\mathbf{p}_S \in \mathcal{S}(p)$

Proposition

The extreme ray densities of $\mathcal{S}_d(p)$ are

$$p_{j_1, j_2}(y) = \begin{cases} \frac{j_2 - pd}{j_2 - j_1} & y = j_1 \\ \frac{pd - j_1}{j_2 - j_1} & y = j_2 \\ 0 & \text{otherwise} \end{cases}, \quad (8)$$

with $j_1 = 0, 1, \dots, j_1^M$, $j_2 = j_2^m, j_2^m + 1, \dots, d$, j_1^M is the largest integer less than pd and j_2^m is the smallest integer greater than pd . If pd is integer the extreme ray densities contain also

$$p_{pd}(y) = \begin{cases} 1 & y = pd \\ 0 & \text{otherwise} \end{cases}. \quad (9)$$

Corollary

- ① If pd is not integer there are $n_p = (j_1^M + 1)(d - j_1^M)$ extreme ray densities.
- ② If pd is integer there are $n_p = d^2 p(1 - p) + 1$ extreme ray densities.

Effect of p on S : bounds

The VaR in $S_d(p)$ is known in closed form.

Proposition-[Luciano, Fontana and Semeraro (2021)]

Let us consider the class $S_d(p)$. Let j_1^M be the largest integer smaller than pd , j_2^m be the smallest integer greater than pd and $j_1^p = \frac{(p-(1-\alpha)d)}{\alpha}$.

- ① If $p < 1 - \alpha$, $\min_{S \in S_d(p)} \text{VaR}_\alpha(S) = 0$ and $\max_{S \in S_d(p)} \text{VaR}_\alpha(S) = \lfloor \frac{pd}{1-\alpha} \rfloor$ if $\frac{pd}{1-\alpha}$ is not integer and $\max_{S \in S_d(p)} \text{VaR}_\alpha(S) = \frac{pd}{1-\alpha} - 1$ if it is integer.
- ② If $1 - \alpha \leq p \leq 1 - \alpha + \frac{\alpha}{d} j_1^M$, $\min_{S \in S_d(p)} \text{VaR}_\alpha(S) = j_1^*$, where j_1^* is the smallest integer greater or equal to j_1^p and $\max_{S \in S_d(p)} \text{VaR}_\alpha(S) = d$.
- ③ If $p > 1 - \alpha + \frac{\alpha}{d} j_1^M$, $\min_{S \in S_d(p)} \text{VaR}_\alpha(S) = j_2^m = j_1^M + 1$ and $\max_{S \in S_d(p)} \text{VaR}_\alpha(S) = d$. In this case, if pd is integer $j_1^M + 1 = pd$.

Let $ES_\alpha(S_d)$ be its expected shortfall. Then

$$\min_{S \in S_d(p)} \text{VaR}_\alpha(S) \leq ES_\alpha(S_d) \leq d.$$

Individual default probability and random covariates

The risk of an homogeneous exchangeable portfolio of obligors is completely determined by the distribution of Q , the random variable defining the common default probability of the homogeneous portfolio. Different estimates of Q lead to different risk valuations.

We assume that the mixing variable Q is a function h of random observable covariates $\mathbf{X} = (X_1, \dots, X_n)$, representing the obligors characteristics. Formally,

$$Q = h(\mathbf{X}).$$

The realizations of Q are functions of the realizations of $\mathbf{X}$, $q = h(\mathbf{x})$.

Using ML techniques, we obtain a sample of observation of Q from a sample of obligors with observable characteristics. Obviously different ML techniques gives different estimates of the sample default probabilities and therefore of the moments and parameters of Q .

Research methodology-in vitro

In-vitro analysis

- 1 ML model risk on individual probabilities: synthetic sample of obligors from random covariates using the logistic regression, then we estimate the different ML models and measure the error in predicting single probabilities and we estimate the α -quantile of the individual probability distribution.
- 2 ML model risk for exchangeable portfolios: we compute the VaR bounds for the class of exchangeable portfolios with marginal default probability (and equi-correlation) estimated using each ML model.
- 3 ML Model risk under a specific model: we calibrate a Bernoulli mixture model using the real synthetic individual probabilities and the estimates obtained using the ML methods. We then compute the VaR and compare the results.

Research methodology-real data

Real data application

- 1 ML model risk on individual probabilities: we estimate the default probabilities using LR and the two ML techniques.
- 2 ML model risk for exchangeable portfolios: we compute the VaR bounds computation for each ML model of the exchangeable portfolio.
- 3 ML Model risk under a specific model: we calibrate a Bernoulli mixture model using the sample individual probabilities obtained with the LML methods. We then compute the VaR and compare the results.

In-vitro analysis - synthetic database

We calibrate the parameters in order to observe 20% of defaults (as in the real dataset Kaggle).

- **uniform_independent**: It is generated through a logit model having two covariates, X_1, X_2 i.i.d. $\mathcal{U}[0, 1]$, i.e., and $\beta = [-0.2, 1.5, -5.0]$,
- **2_squared**: It is generated through a logit model having two covariates X_1, X_2 , where $X_1 \sim \mathcal{U}[0, 1]$ and $X_2 = X_1^2 + U[-0.05, 0.05]$ and $\beta = [-0.2, 0.7, -5.5]$,
- **normal_copula**: Logit model with two covariates: marginals are $\mathcal{U}[0, 1]$ linked through a Normal copula which covariance matrix is

$$\begin{bmatrix} 0.5 & -0.2 \\ -0.2 & 0.5 \end{bmatrix}. \quad (10)$$

Moreover, we set $\beta = [-0.2, 0.5, -3.2]$,

- **t_copula**: Logit model with the two marginals linked through a t copula with 2 degrees of freedom and $\beta = [-0.2, 0.5, -3.2]$,
- **5_non_linear**: Logit model where the covariates are: X_1, X_2, X_3 i.i.d. $\mathcal{U}[0, 1]$ $X_1 \cdot X_2 + \mathcal{U}[-0.05, 0.05]$, $X_1 \cdot X_3 + \mathcal{U}[-0.05, 0.1]$, and $\beta = [-0.1, -1, -0.5, -0.5, -1, -1]$,

The choice of the uniform distribution between 0 and 1 represents the values of the features after an opportune scaling.

In-vitro analysis - imbalance

For each of the settings we generate 500, 1000, 5000, and 10000 records, and 5 datasets for each possible choice, resulting in a total of $5 \times 4 \times 5 = 100$ datasets. We try to apply different techniques for dealing with imbalanced datasets:

Imbalance method	Average error [%]
ClusterCentroids	24.14 (3.36)
Identity (no imbalance)	06.40 (5.01)
RandomOverSampler	24.50 (3.39)
RandomUnderSampler	24.81 (3.21)
SMOTE	24.12 (3.22)
SMOTETomek	23.85 (3.35)
TomekLinks	07.18 (4.65)

Table: Average absolute error on p for different imbalance techniques.

Despite the imbalance techniques are of paramount importance for the classification problem, they do not improve the performance of the methods for the default probability estimation.

We continue the experiments by not considering any imbalance algorithm.

In-vitro analysis - results

The best results are obtained by the LogisticRegression and by the MLPClassifier: these two methods do not have model errors.
The third best method is the Random Forest Classifier.

	uniform_ind	2_squared	n_copula	t_copula	5_n_1
KNN	11.25(0.80)	12.63(0.89)	13.76(0.84)	13.76(1.17)	13.67(0.66)
LogisticRegression	2.05(1.61)	2.27(1.46)	1.91(1.62)	1.69(1.09)	1.43(0.96)
MLPClassifier	1.50(1.25)	2.53(1.54)	2.05(1.63)	1.66(0.99)	1.66(1.14)
RandomForestClassifier	6.51(0.63)	3.16(1.11)	3.93(0.94)	4.62(0.33)	3.10(0.76)

Table: Average percentage error for different ML techniques.

Due to the bad results, in the following, we do not consider KNeighborsClassifier.

In-vitro analysis - results

id_setting	n_data	p	rho	p1_LR	rho_LR	p1_RF	rho_RF
Uniform	500	0.18	0.22	0.19	0.14	0.20	0.13
	1000	0.19	0.24	0.21	0.19	0.20	0.10
	5000	0.20	0.23	0.21	0.21	0.21	0.09
	10000	0.20	0.24	0.21	0.24	0.21	0.11
Normal	500	0.18	0.07	0.23	0.04	0.22	0.04
	1000	0.21	0.09	0.24	0.12	0.23	0.10
	5000	0.20	0.09	0.20	0.09	0.20	0.05
	10000	0.20	0.09	0.21	0.08	0.21	0.04
5_non_lin	500	0.20	0.06	0.19	0.03	0.19	0.04
	1000	0.19	0.05	0.20	0.06	0.20	0.04
	5000	0.19	0.06	0.18	0.06	0.18	0.04
	10000	0.19	0.06	0.19	0.07	0.18	0.04

Table: Estimated moments

In-vitro analysis - results

id_setting	n_data	a	b	a_LR	b_LR	a_RF	b_RF
Uniform	500	0.61	2.85	1.16	4.95	1.32	5.37
	1000	0.62	2.60	0.91	3.45	1.75	6.84
	5000	0.67	2.60	0.80	3.06	2.03	7.74
	10000	0.63	2.50	0.64	2.45	1.67	6.38
Normal	500	2.30	10.43	5.17	17.73	4.86	17.14
	1000	2.07	7.57	1.80	5.79	2.13	7.22
	5000	2.04	7.94	2.02	7.88	3.76	14.61
	10000	2.14	8.37	2.33	9.04	4.49	17.29
5_non_lin	500	3.36	13.44	5.69	23.99	5.16	22.28
	1000	3.38	14.16	3.34	13.72	4.36	17.41
	5000	2.90	12.58	2.74	12.08	4.11	18.23
	10000	2.90	12.73	2.54	11.13	4.83	21.30

Table: Beta parameters

In-vitro analysis - results

id_setting	BetaTail	BetaTailLR	BetaTailRF
Uniform	112.00	118.00	117
	218.00	235.00	228
	1053.00	1073.00	1073
	2081.00	2122.00	2133
Normal	113.00	134.00	133
	218.00	242.00	236
	1048.00	1045.00	1048
	2074.00	2096.00	2104
5_non_lin	104.00	101.00	101
	202.00	205.00	207
	966.00	949.00	951
	1914.00	1914.00	1914

Table: Beta tail: 0.9-quantile

In-vitro analysis - results

We do not report the upper bound since it is equal to N .

Setting	n_data	minVar	BetaBinVar	BetaTail
Uniform	1000	104	485	218
	10000	1135	5015	2081
Normal	1000	127	391	218
	10000	1151	3689	2074
5_non_lin	1000	103	319	202
	10000	949	3164	1914

Table: Synthetic Portfolio VaR

Setting	n_data	minVarLR	BetaBinVarLR	minVarRF	BetaVarRF
Uniform	1000	119	466	115	388
	10000	1183	5104	1193	3975
Normal	1000	152	443	141	410
	10000	1169	3646	1178	3209
5_non_lin	1000	106	325	111	315
	10000	951	3258	944	2858

Table: Estimated VaR

Real data application

Individual probabilities and parameter estimates

For each choice of $h = \text{LR, RF, AB, KNN}$.

- 1 We estimate $q_i^h = h(\mathbf{x}_i)$, where $\mathbf{x}_i$ is the m -dimensional vector of covariates realizations, and find a sample of estimated conditional default probabilities $\hat{\mathbf{q}}^h = (\hat{q}_1^h, \dots, \hat{q}_n^h)$;
- 2 we compute the marginal default probability p and the equicorrelation among default indicators.
- 3 we estimate the parametrical beta-binomial distribution of the number of default S by moments matching, using the first two moments of Q_h , i.e. p and π_2 ;

Real data application

- 1 we analyse the effect of p on the VaR, by using analytical bounds;
- 2 we analyze the effect of the first two moments p and π_2 on the risk of the aggregate loss, by computing the VaR and ES of the beta-binomial distribution of the loss;

Real data application Data

Data description

- Kaggle database on credit card defaulters in Taiwan, from April to September 2005. Sample size of $n = 30'000$ obligors and $m = 24$ covariates.
- The covariates consist of age, marital status, monthly repayment status, past bill amounts and others.
- From the label frequencies in the outcome variable **Y**, we see that the dataset is **unbalanced**, but since the unbalancing is not extreme, we did not correct it.

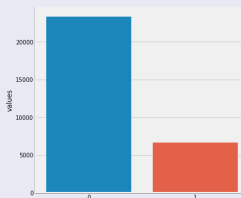


Figure: Unbalancing of the dataset. Default=1, non-default=0.

Results- Individual probabilities and parameter estimates

ML techniques and Evaluation metrics

- **Machine Learning techniques** involved:
 - Random Forest
 - AdaBoost
 - K-Nearest Neighbours
- Their performances are compared against the **Logistic Regression** model. We consider four different evaluation metrics to do the comparison, giving priority to the **F1-score metric** to select the best model.

Model	Precision	Recall	F1-score	AUC
LR	0.61	0.78	0.68	0.64
RF	0.79	0.81	0.78	0.77
AB	0.80	0.82	0.79	0.77
KNN	0.78	0.81	0.78	0.72

Table: Performance measure for each model.

Individual probabilities and parameter estimates

ML techniques and Evaluation metrics

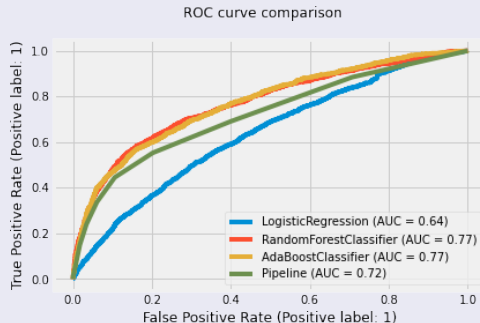


Figure: Overlapped ROC curve for each model specification, with relative AUC score.

Results- model parameters

The results on the real datasets are the following:

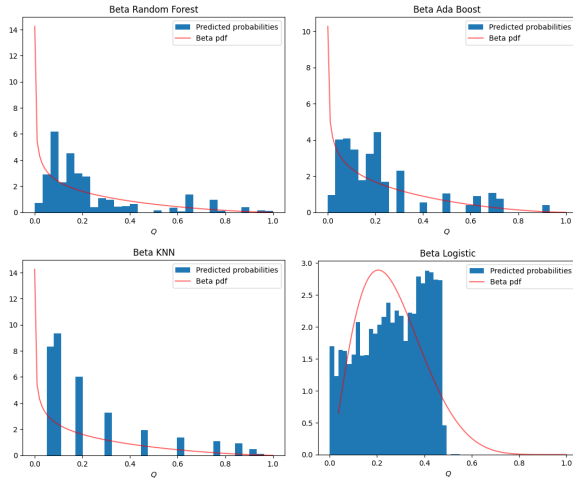
	Beta parameters			
	a	b	p	rho
LR	2.42	6.77	0.26	0.09
RF	0.69	2.41	0.22	0.24
AB	0.73	2.57	0.22	0.23
KNN	0.69	2.41	0.22	0.24

Table: Beta parameters, individual default probability and portfolio equicorrelation

- LR **overestimates** the marginal default probability and **underestimates** the second order moment, if compared with ML methods.
- **Default correlation** is higher for ML models because they incorporate **linear** and **non linear** dependencies among covariates.

Results - distribution of individual default probability

Figure: Empirical vs beta distribution of individual defaults.



Results-Portfolio risk

	VaR_α			Bound_α		
α :	0.9	0.95	0.99	0.9	0.95	0.99
LR	2731	3110	3796	(1089, 6000)	(1347, 6000)	(1535, 6000)
RF	3205	3870	4885	(818, 6000)	(1091, 6000)	(1289, 6000)
AB	3155	3801	4806	(822, 6000)	(1095, 6000)	(1293, 6000)
KNN	3205	3870	4885	(818, 6000)	(1091, 6000)	(1289, 6000)

Table: Beta-binomial Var and Bounds

ML effect on exchangeable portfolios: the minimum VaR is slightly higher for LR, while all the ML techniques perform similarly.

ML effect on the beta-binomial model: the VaR for the ML methods is similar and higher if compared to LR.

ML estimates higher correlations: correlation affects the risk of the portfolio more than individual default probability

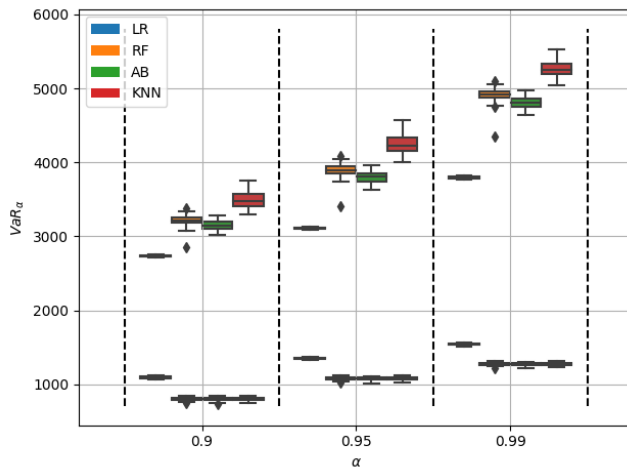


Figure: Beta-binomial and lower bound distributions across different training sets

Ongoing research: the effect of ρ .

Let us consider the class $\mathcal{S}_d(p, \rho)$.

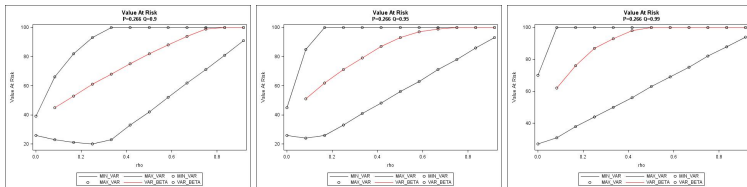
- 1 Ray densities and their VaR_α are analytical: bounds are found by computationally searching the maximum and the minimum VaR_α among ray densities.
- 2 We deal with $\mathcal{E}(p, \rho)$ in the three scenarios $p = 0.3\%$, $p = 1.7\%$ and $p = 26.6\%$ and provide bounds for VaR_α for three levels of correlation: $\rho = \frac{1}{6}; \frac{1}{2}; \frac{5}{6}$.
- 3 We also report the VaR corresponding to an exchangeable Bernoulli mixing model from the credit risk literature: β -mixing model.

The effect of ρ :VaR bounds

The class $\mathcal{E}_d(p, \rho)$

We have similar results on the subclass $\mathcal{S}_d(p, \rho)$ and therefore on the subclass $\mathcal{E}_d(p, \rho)$: analytical ray densities and bounds for VaR.

Figure: VAR bounds and β -mixing model VAR for $p = 26.6\%$, $d = 100$ and different ρ s



Ongoing research: portfolio risk measures

Expected shortfall

Let S be a random variable representing the number of defaults with finite mean. The $\mathbf{ES}_\alpha$ at level α is defined by

$$\mathbf{ES}_\alpha(S) = \frac{1}{1-\alpha} (E[S; S \geq \text{VaR}(S)] + \text{VaR}(S)(1 - \alpha - P(S \geq \text{VaR}(S)))).$$

Expected shortfall is a convex measure of risk

ES bounds

Fontana and Semeraro (2020)

The minimum convex sums is on $S_{m,m+1}$, that is the extremal random variable in $S_p(dp)$ with support on the two integer around the mean pd , say $m, m + 1$.
Th maximum is on the Upper Fréchet bound

Consequence

We can explicitly find sharp bounds for the ES, since it is a convex measure of risk.

Preliminary results - ongoing

ES Estimates- β Binomial model

We perform a preliminary analysis in Doria, Luciano and Semeraro on an exchangeable portfolio comparing LR and AB.

α	ES LR	ES AB
0.99%	4099.6	5131.9
0.95%	3524.3	4395.1
0.90%	3213.0	3916.9

Table: Beta-binomial ES for large portfolios.

Machine Learning techniques in joint default assessment

Conclusion

- Our main result is that **ML methods** can model significantly **higher default correlations** if compared with our benchmark LR,
- Default correlation are the main factors influencing the VaR of the loss, and higher correlations lead to **heavier tails of the loss**.

Ongoing and further research

- Find the VaR bounds for p and ρ estimated with different ML methods
- Find the bounds and the estimated ES for different ML methods
- Consider partially exchangeable portfolios.
- Consider more sophisticated ML techniques.

Thank you,

patrizia.semeraro@polito.it

- Fontana, Roberto, Elisa Luciano, and Patrizia Semeraro. "Model risk in credit risk." *Mathematical Finance* 31.1 (2021): 176-202.
- M. Doria, E. Luciano and P. Semeraro. "Machine Learning techniques in joint default assessment," <https://arxiv.org/abs/2205.01524>, 3 May 2022.