

Applying the Pre-Commitment Approach to bottom up stress tests

by Simone Casellina, Giuseppe Pandolfo, and Mario Quagliariello

Discussion by Peter Raupach, Deutsche Bundesbank

2019 EBA Policy Research Workshop, Paris, 27 November 2019

Summary

– Goals:

- Incentivize banks (bank managers...) to report results of bottom-up stress tests to supervisors *truthfully*.
- Allow banks to run their own risk models, without thorough scrutiny.

– Setup of bottom-up stress tests:

- Regulator defines stress test scenario(s).
- Bank runs own model to forecast key figures conditional on the scenario. Bank could use a model $\mathbf{P}^{report}$ different from its “best belief” $\mathbf{P}^{honest}$.

– Devices:

- **Self-revelation**: Bank reports its forecast future loss rate (or related).
- An ex-post **charge** paid by the bank(er), equivalent to the mean squared error in regression analysis

$$charge = \delta \left(LR_{t+h} - \mathbf{E}_t^{report}(LR_{t+h}) \right)^2$$

- Important: $\mathbf{E}_t^{report}(LR_{t+h})$ is not the stress testing result but the „plain“ unconditional forecast of LR_{t+h} .

A major concern (I think...)

- The charge would actually incentivize truthful reporting of the unconditional forecast: the bank would choose a reporting model such that

$$\mathbf{E}_t^{report}(LR_{t+h}) = \mathbf{E}_t^{honest}(LR_{t+h}) \quad (H)$$

because this minimizes $\mathbf{E}_t^{honest} \left(LR_{t+h} - \mathbf{E}_t^{report}(LR_{t+h}) \right)^2$

- Does it imply $\mathbf{P}^{report} = \mathbf{P}^{honest}$? No.
- What are we actually interested in? Truthful stress testing results:

$$\mathbf{E}_t^{report}(LR_{t+h}|STScen) = \mathbf{E}_t^{honest}(LR_{t+h}|STScen)$$

- Can we backtest/incentivize this? No, stress scenarios are rare.
- ➔ **Underlying assumption:** A manipulated model would hurt (H) and hence raise the expected charge.
- Can we hope for that? ➔ See next slide.

Example from the paper's appendix

- “Agreed” macro model: GDP growth $y_t \sim AR(1)$, stress scenario $\{y_{t+1} = y^*\}$
- Default rate model $\mathbf{P}^{honest}$: $DR_{t+1} = d + \rho DR_t + \theta y_{t+1} + \epsilon_{t+1}, \quad \theta < 0$

- **Cheating strategy** (in the paper): constant bias, i.e.

$$\mathbf{P}^{report}: DR_{t+1}^{rep} = d + \rho DR_t + \theta y_{t+1} + \epsilon_{t+1} - \mathbf{b}$$

- This shows up in the charge! → Incentive device works.

- **Another cheating strategy**: Choose $C > y^*$ and limit the impact of bad y_{t+1}

$$\mathbf{P}^{report}: DR_{t+1}^{rep} = d + \rho DR_t + \theta \max(y_{t+1}, C) + \epsilon_{t+1}$$

- Most of the time, the model is honest:

$$\mathbf{E}_t^{report}(DR_{t+1} | y_{t+1} > C) = \mathbf{E}_t^{honest}(DR_{t+1} | y_{t+1} > C)$$

- But not conditional on stress:

$$\mathbf{E}_t^{report}(DR_{t+1} | y_{t+1} = y^*) = \mathbf{E}_t^{honest}(DR_{t+1} | y_{t+1} = C) < \mathbf{E}_t^{honest}(DR_{t+1} | y_{t+1} = y^*)$$

- **Incentive device is blind to “tail cosmetics”**

Conclusion

- The incentive device is **not reliable** if banks can distort the model's **relationship** between frequently observed events and rarely observed scenarios (of interest in stress tests)
- Two ways out:
 - **A**: Thorough scrutiny of bank's own models;
 - **B**: Prescribe the model (at least, the **relationship** mentioned above).
- Designers of the Basel 2 IRB approach chose option **B**:
 - Let banks determine PDs (+ LGDs, EADs in AIRBA), require PD backtests.
 - Impose the Vasicek model on all banks to determine loss tails.

Minor Comments

- Charge is **symmetric to under- and overestimating** the effect of stress
→ targeted at “being right as often as possible”.
- But these errors have **different costs** and, more importantly, different externalities. Intuitively, underestimation is worse.
- Compare with **credit scoring** tools, virtually all standard **medical tests**: they rarely make alpha errors but frequently beta errors.
- Think about **tilted charges** derived from costs. General framework:
 - Gneiting, T., & Raftery, A. E. (2007). Strictly proper scoring rules, prediction, and estimation. *JASA* 102(477), 359-378.
- Relevant reference not in the paper:
 - Migueis, M. (2017). Forward-looking and incentive-compatible operational risk capital framework. *SSRN* 2964945.