

A method to find approximate solutions of first order systems of non-linear ordinary equations.

J.J. Alvarez-Sánchez¹, M. Gadella², L.P. Lara^{3,4}

¹ Escuela de Informática de Segovia, Universidad de Valladolid, Segovia, Spain

² Departamento de Física Teórica, Atómica y Optica and IMUVA,
Universidad de Valladolid, 47011 Valladolid, Spain

³ Instituto de Física Rosario, CONICET-UNR,
Bv. 27 de Febrero, S2000EKF Rosario, Santa Fe, Argentina.

⁴ Departamento de Sistemas, Universidad del Centro Educativo
Latinoamericano, Av. Pellegrini 1332, S2000 Rosario, Argentina

March 4, 2021

Abstract

We develop a one step matrix method in order to obtain approximate solutions of first order non-linear systems and non-linear ordinary differential equations, reducible to first order systems. We find a sequence of such solutions that converge to the exact solution. We apply the method to different well known examples and check its precision, in terms of local error, comparing it with the error produced by other methods. The advantage of the method over others widely used lies on the great simplicity of its implementation.

1 Introduction

This paper pretends to be a contribution to methods to find the approximate solutions of nonlinear first order equations (or systems) with given initial values of the form

$$\mathbf{y}'(t) = \mathbf{f}(\mathbf{y}(t)), \quad \mathbf{y}(t_0) = \mathbf{y}_0, \quad (1)$$

where the prime represents first derivative with respect to the variable t . As many higher order ordinary differential equations either linear or non-linear may be written as a system of the form (1), our method to obtain approximate solutions will apply also to these kind of equations.

Our approach is based in a generalization of the one-step matrix method developed by Demidovich and other authors [1–3] some time ago, valid for systems of the form $\mathbf{y}'(t) = A(t)\mathbf{y}(t)$. One clear presentation of this method appears in the textbook by Farkas [4] and it is interesting to compare it with some related procedures, see for instance [6, 7]. It is quite important to remark that, while the Demidovich matrix method is applied to linear equations

(with variable coefficients), our matrix method is a generalization to non-linear systems. The advantage that our method may have in comparison of other one step methods lies on its great algorithmic simplicity. In addition, our solutions have a reasonable level of accuracy in few steps and this is workable in a table computer. This method is quite easily programmable and is very suitable for its use with the package Mathematica. It may be also seen as an alternative to Runge-Kutta and Taylor methods due precisely to its simplicity and precision. The objective of the present study is to obtain approximate solutions of all kind of systems described by non-linear systems of differential equations (although we may include those linear systems with variable coefficients), including those who appear from physics.

In the derivation of the present approach, our motivation was rooted in the practice of operational numerical calculus. Nevertheless, we are mainly focused in the mathematical analysis of our method instead of a detailed analysis of the algorithm or CPU times. As for the case of one step Taylor polynomial method, ours shows a great conceptual simplicity. As a consequence, our proposal for obtaining approximate solutions of non-linear systems can be very easily implemented.

The presentation of the method to obtain the approximate solutions is introduced in Section 2. Thus, we have a sequence of approximate solutions, which are defined on a given interval of the real line. This sequence converges uniformly to the exact solution on the given interval as is proven in Section 3, where we also discuss a question of order. It is important to remark that we do not impose any periodicity conditions, so that our results are valid either for periodic or for non-periodic solutions.

We have applied our method to various examples of two or three dimensional examples. It is also necessary to test the applicability and accuracy of the method on widely used equations and/or systems. To this end, we have used the van der Pol [5], Duffing [8], Lorentz [9] equations, a pseudo diffusive equation depending on a parameter, studied in the standard literature [10], an epidemic equation and the predator-prey Lotka-Volterra [11, 12] equation. We have compared the precision of our solutions with the exact solution, whenever this is known. If not, the comparison is based on Runge-Kutta solutions, for a modern presentation, see [13–15]. Also with the widely used Taylor method.

The analysis of a variety of examples suggest that our method may be more precise for two dimensional systems than for higher dimensional ones, although this is not always exact: we give two examples of three dimensional systems (Lorentz and epidemic equations), for which our method gives different precisions. And it looks like particularly efficient in the case of the pseudo diffusive equation we mentioned earlier, at least when compared with the standard perturbative method studied in [16].

We usually obtain precisions in between of those obtained for the Taylor method of third and fourth order (although the precision depends also on the length of subintervals in which is divided the interval in which we are looking for solutions), which is reasonable when we use a table computer. We close this article with concluding remarks and a conjecture on the Liénard equation which is motivated by some of our results and confirmed through numerical experiments.

2 A matrix method.

We begin with an initial value problem as given in (1). While the function $\mathbf{y}(t)$ is a $\mathbb{R}^n$ valued function with real variable t of class C^1 on the neighbourhood $|t - t_0| \leq a$, $\mathbf{f}(-)$ is a $\mathbb{R}^n$ valued function with variable in $\mathbb{R}^n$, which is continuous on $D \equiv \{\mathbf{y} \in \mathbb{R}^n \mid \|\mathbf{y} - \mathbf{y}_0\| \leq d\}$ and satisfies a Lipschitz condition with respect to $\mathbf{y}$. Needless to say that a, d are positive constants.

In relation with the identity (1), we shall use either the denomination of “equation” or “system” indistinctly. In any case, it is well known that the initial value problem (1) has one unique solution on the interval $|t - t_0| \leq \inf(a, d/M)$ with $M := \sup_D \|\mathbf{f}\|$.

Our objective is to introduce a generalization of a method of solutions of (1) proposed in [4]. This generalization is based in an iterative procedure of numerical integration for equations of the type $\mathbf{f}(\mathbf{y}) = A(\mathbf{y}) \cdot \mathbf{y}$, where $A(-)$ is an $n \times n$ matrix. With this choice for the function $\mathbf{f}(-)$, the differential equation in (1) has the form

$$\mathbf{y}'(t) = A(\mathbf{y}(t)) \cdot \mathbf{y}(t). \quad (2)$$

We assume that the entries of $A(\mathbf{y})$ are continuous on D . Let us define a uniform partition of the interval $[0, t_N]$ into subintervals $I_k \equiv [t_k, t_{k+1}]$, where $t_k = hk$ with $k = 0, 1, 2, \dots, N$, with N natural and $h = t_N/N$, $t_N < a$. We have chosen this form of the interval by simplicity, needless to say that if the original interval is somehow else, it always may be transformed into $[0, t_N]$ by a translation. Also, the equal spacing of all subintervals is not strictly necessary, although it simplifies our notation. Conventionally, we may call *nodes* to the points $\{t_k\}$.

We proceed as follows: On each interval I_k , we approximate Equation (2) by

$$\mathbf{y}'_{N,k}(t) = A(\mathbf{y}_{N,k}^*) \cdot \mathbf{y}_{N,k}(t). \quad (3)$$

At each node, t_k , we impose $\mathbf{y}_{N,k}(t_k) = \mathbf{y}_{N,k-1}(t_k)$, while $\mathbf{y}_{N,k}^* \in \mathbb{R}^n$ is a constant to be determined. This gives the segmentary solution, which has to be of the form $\mathbf{y}_N(t) \equiv \{\mathbf{y}_{N,k}(t) \mid k = 0, 1, \dots, N-1\}$. It satisfies

$$\mathbf{y}'_N(t) = A(\mathbf{y}_N^*) \cdot \mathbf{y}_N(t), \quad (4)$$

where $\mathbf{y}_N^*$ coincides with $\mathbf{y}_{N,k}^*$ on each of the k -th intervals.

These segmentary solutions give a sequence of functions $\{\mathbf{y}_N(t)\}$, $t \in [0, t_N]$, that are approximations to the solution of (2). Here, we shall give a method to obtain each of the $\mathbf{y}_N(t)$ and, in next sections, we shall discuss the properties of the sequence.

Then, we proceed to an iterative integration of (4) as follows: Take the first interval I_0 and fix an initial value $\mathbf{y}(t_0)$. This initial value gives the solution $\mathbf{y}_{N,0}(t)$ on I_0 . Thus, we have the value $\mathbf{y}_{N,0}(t_1) = \mathbf{y}_{N,1}(t_1)$, which serves as the initial value for the solution on the interval I_1 . Then, we repeat the procedure in an obvious manner for I_2 and so on. For each interval I_k the matrix $A(\mathbf{y}_{N,k}^*)$, which appears in (3), is a constant matrix and, therefore, (3) is a system with constant coefficients. Therefore, the solution of (3) has the form

$$\mathbf{y}_{N,k}(t) = \exp\{A(\mathbf{y}_{N,k}^*)(t - t_k)\} \cdot \mathbf{y}_{N,k}(t_k), \quad k = 0, 1, 2, \dots, N-1, \quad (5)$$

where we have used the notation $\mathbf{y}_{N,0}(t_0) = \mathbf{y}(t_0)$. We determine the numbers $\mathbf{y}_{N,k}^*$ through the following expression:

$$\mathbf{y}_{N,k}^* = \mathbf{y}_{N,k} \left(t_k + \frac{h}{2} \right) = \exp \left\{ A(\mathbf{y}_{N,k}(t_k)) \frac{h}{2} \right\} \cdot \mathbf{y}_{N,k}(t_k), \quad (6)$$

where $k = 0, 1, 2, \dots, N-1$.

Then, the approximate solution $\mathbf{y}_N(t)$ gives at $t \in I_k$ and on each of the intervals I_k :

$$\mathbf{y}_N(t) = \exp \left\{ A(\mathbf{y}_{N,k}^*)(t - t_k) \right\} \cdot \prod_{j=0}^{k-1} \exp \left\{ A(\mathbf{y}_{N,j}^*) h \right\} \cdot \mathbf{y}_0. \quad (7)$$

The determination of the exponential of a matrix may often be rather complicated for large matrices. Then, we may use the Putzer spectral formula. This establish that if A is a constant matrix of order $n \times n$ with eigenvalues $\{\lambda_k\}_{k=1}^n$, then, its exponential verifies the following expression:

$$\exp\{A t\} = \sum_{k=1}^n r_k(t) P_{k-1}, \quad (8)$$

with

$$P_0 \equiv I, \quad P_k = \prod_{j=1}^k (A - \lambda_j I), \quad k = 1, 2, \dots, n-1, \quad (9)$$

where I is the $n \times n$ identity matrix and the coefficients $r_k(t)$ are to be determined through the following first order system of differential equations:

$$\begin{aligned} r_1'(t) &= \lambda_1 r_1(t), \quad r_1(0) = 1; \\ r_k'(t) &= \lambda_k r_k(t) + r_{k-1}(t), \quad r_k(0) = 0, \end{aligned} \quad (10)$$

for $k = 2, 3, \dots, n$.

For simplicity, let us consider the particular case, in which $A(\mathbf{y}_{N,j}^*)$ are matrices of order 2×2 . Each of these matrices has two eigenvalues, $\lambda_{1,j}$ and $\lambda_{2,j}$, which may be either different or equal. Let us assume that $\lambda_{1,j} \neq \lambda_{2,j}$. Then,

$$\begin{aligned} &\exp \left\{ A(\mathbf{y}_{N,j}^*) h \right\} \\ &= \frac{1}{\lambda_{1,j} - \lambda_{2,j}} \left\{ (A(\mathbf{y}_{N,j}^*) - \lambda_{2,j} I) \exp\{\lambda_{1,j} h\} - (A(\mathbf{y}_{N,j}^*) - \lambda_{1,j} I) \exp\{\lambda_{2,j} h\} \right\}. \end{aligned} \quad (11)$$

On the other hand, when $\lambda_{1,j} = \lambda_{2,j} = \lambda_j$, we have for the exponential

$$\exp \left\{ A(\mathbf{y}_{N,j}^*) h \right\} = \exp\{\lambda_j h\} \{I + h(A(\mathbf{y}_{N,j}^*) - \lambda_j I)\}. \quad (12)$$

We conclude here the crude description of the method. In the sequel, we shall show the convergence of the sequence, $\{\mathbf{y}_N(t)\}$, of approximate segmentary solutions and shall evaluate the precision of the method.

3 On the convergence of approximate solutions

In the previous Section, we have obtained a set of approximate solutions of the initial value problem on a compact interval of the real line. The question is now, assuming we have obtained by the previous method a sequence of solutions. Does this sequence converges to the exact solution in any reasonable sense as the length of the sub intervals, here called h , becomes arbitrarily small. To investigate this possibility is the goal of the present Section. Let us go back to (4) and rewrite it as

$$\mathbf{y}'_N(t) = A(\mathbf{y}_N) \cdot \mathbf{y}_N(t) + \eta_N, \quad (13)$$

so that,

$$\eta_N(t) = (A(\mathbf{y}_N^*) - A(\mathbf{y}_N)) \cdot \mathbf{y}_N(t). \quad (14)$$

Let us add and subtract $A(\mathbf{y}_N^*) \cdot \mathbf{y}_N^*$ in the right hand side of (14). Then, let us take the supremum norm on the interval $[t - t_0, t + t_0]$ and use the triangle inequality of the norm, so as to obtain the following inequality:

$$\|\eta_N\| \leq \|(A(\mathbf{y}_N^*) \cdot \mathbf{y}_N^* - A(\mathbf{y}_N) \cdot \mathbf{y}_N)\| + \|A(\mathbf{y}_N^*)\| \|\mathbf{y}_N^* - \mathbf{y}_N\|. \quad (15)$$

Then, we apply in (15) the Lipschitz condition with respect to the variable $\mathbf{y}$ with constant $K > 0$. Then, it comes that

$$\|\eta_N\| \leq (K + \|A(\mathbf{y}_N^*)\|) \|\mathbf{y}_N^* - \mathbf{y}_N\|. \quad (16)$$

On each one of the intervals I_k , let us expand $\mathbf{y}_N(t)$ in Taylor series around t_k . We obtain the following inequality:

$$\|\mathbf{y}_N^* - \mathbf{y}_N(t)\| \leq \frac{h}{2} \|\mathbf{y}'_N(t_k)\| \leq \frac{h}{2} \max_{t_0 \leq t \leq t_N} \|\mathbf{y}'_N(t)\|. \quad (17)$$

Since $\mathbf{y}'_N(t)$ is continuous on the interval $t_0 \leq t \leq t_N$, the maximum in the right hand side of (17) exists. Furthermore, $A(\mathbf{y})$ is continuous with respect to $\mathbf{y}$ on the neighborhood $\|\mathbf{y} - \mathbf{y}_0\| \leq d$, so that there exists a constant $C > 0$ such that $\|A(\mathbf{y})\| \leq C$. Taking norms in (4), we have that

$$\|\mathbf{y}'_N\| \leq \|A(\mathbf{y}_N^*)\| \|\mathbf{y}_N\| \leq C \|\mathbf{y}_N\|. \quad (18)$$

Equation (7) implies that

$$\max_{t_0 \leq t \leq t_N} \|\mathbf{y}_N(t)\| \leq C' \exp\{C(t_N - t_0)\}, \quad (19)$$

where $C' > 0$ is a constant. After (18-19), we see that $\|\mathbf{y}'(t)\|$ is bounded for all t in the interval $[t_0, t_N]$. Consequently after (16) and (18-19), we have that

$$\|\eta_N\| \leq \frac{h}{2} (K + C) C' \exp\{C(t_N - t_0)\} = S h, \quad (20)$$

where the meaning of the constant $S > 0$ is obvious.

Next, let us integrate (13) on the interval $[t_0, t]$. Since for all value of N , we use the same initial value $y(t_0)$, we have

$$\mathbf{y}_N(t) = \mathbf{y}(t_0) + \int_{t_0}^t A(\mathbf{y}_N(s)) \cdot \mathbf{y}_N(s) ds + \int_{t_0}^t \eta_N(s) ds. \quad (21)$$

From (21), we have that

$$\|\mathbf{y}_{N+M}(t) - \mathbf{y}_N(t)\| \leq \int_{t_0}^t \|A(\mathbf{y}_{N+M}(s)) \cdot \mathbf{y}_{N+M}(s) - A(\mathbf{y}_N(s)) \cdot \mathbf{y}_N(s)\| ds + \delta_N(t), \quad (22)$$

with

$$\delta_N(t) = \int_{t_0}^t \|\eta_{N+M}(s) - \eta_N(s)\| ds \leq 2Sh(t_N - t_0). \quad (23)$$

Using the Lipschitz condition in (22), we obtain

$$\|\mathbf{y}_{N+M}(t) - \mathbf{y}_N(t)\| \leq K \int_{t_0}^t \|\mathbf{y}_{N+M}(s) - \mathbf{y}_N(s)\| ds + 2Sh(t_N - t_0). \quad (24)$$

At this point, we use the Gronwall lemma, which states the following

Lemma (Gronwall).- Let $f(t) : I \mapsto \mathbb{R}$ an integrable function on the compact real interval I such that there exists two positive constants A and B , with

$$0 \leq f(t) \leq A + B \int_{t_0}^t f(s) ds, \quad t_0 \in I \quad (25)$$

for all $t \in I$. Then,

$$f(t) \leq A e^{B(t-t_0)}. \quad (26)$$

■

Then, we use the Gronwall lemma with

$$f(t) \equiv \|\mathbf{y}_{N+M}(t) - \mathbf{y}_N(t)\|, \quad A \equiv Mh := 2S(t_N - t_0)h, \quad B \equiv K, \quad (27)$$

to conclude that

$$\|\mathbf{y}_{N+M}(t) - \mathbf{y}_N(t)\| \leq Mh e^{K(t-t_0)} \leq [M e^{K(t_N-t_0)}] h = K' h, \quad (28)$$

With $K' > 0$ a positive constant. Therefore,

$$\|\mathbf{y}_{N+M}(t) - \mathbf{y}_N(t)\| \mapsto 0, \quad (29)$$

as $h \mapsto 0$. Since the space $C^0[t_0, t_N]$ is complete¹, (29) implies the existence of a continuous function $\mathbf{z}(t) : [t_0, t_N] \mapsto \mathbb{R}$, such that

¹Note that t_N is fixed and N just denotes the number of intervals in the partition or equivalently, the length of h .

$$\mathbf{z}(t) := \lim_{N \rightarrow \infty} \mathbf{y}_N(t), \quad (30)$$

uniformly.

Now, we claim that $\mathbf{z}(t)$ is differentiable on (t_0, t_N) . Furthermore, $\mathbf{z}(t)$ is a solution of the differential equation (2).

The proof goes as follows: The Lipschitz condition applied to our situation implies that

$$\|A(\mathbf{y}_{N+M}(t)) \cdot \mathbf{y}_{N+M}(t) - A(\mathbf{y}_N(t)) \cdot \mathbf{y}_N(t)\| \leq K \|\mathbf{y}_{N+M}(t) - \mathbf{y}_N(t)\|, \quad (31)$$

so that $A(\mathbf{y}_N(t)) \cdot \mathbf{y}_N(t)$ converges uniformly to $A(\mathbf{z}(t)) \cdot \mathbf{z}(t)$. In addition, after (20), we have that

$$\|\eta_N(t)\| \leq Sh \leq S(t_N - t_0). \quad (32)$$

Recall that t_N is fixed for whatever value of N . Then, taken limits in (21), we have

$$\begin{aligned} \mathbf{z}(t) &= \mathbf{y}(t_0) + \lim_{N \rightarrow \infty} \int_{t_0}^t A(\mathbf{y}_N(s)) \cdot \mathbf{y}_N(s) ds + \lim_{N \rightarrow \infty} \int_{t_0}^t \eta_N(s) ds \\ &= \mathbf{y}(t_0) + \int_{t_0}^t \left[\lim_{N \rightarrow \infty} A(\mathbf{y}_N(s)) \cdot \mathbf{y}_N(s) \right] ds + \int_{t_0}^t \lim_{N \rightarrow \infty} [\eta_N(s)] ds. \end{aligned} \quad (33)$$

In the second integral, we may interchange the limit and the integral due to the uniform convergence of the sequence under the integral to its limit. In the case of the second integral, we have used the Lebesgue convergence theorem, which can be applied here due to (32). Since obviously $\lim_{N \rightarrow \infty} [\eta_N(s)] = 0$, we finally conclude that

$$\mathbf{z}(t) = \mathbf{y}(t_0) + \int_{t_0}^t A(\mathbf{z}(s)) \cdot \mathbf{z}(s) ds. \quad (34)$$

From (34), we conclude the following:

- 1.- The function $\mathbf{z}(t)$ is differentiable in the considered interval.
- 2.- The function $\mathbf{z}(t)$ is the solution of equation (2) with initial value $\mathbf{z}(t_0) = \mathbf{y}(t_0)$.

3.1 A question of order

The expansion into Taylor series on a neighborhood of t_k of the solutions of equations (2) and (3) have these forms, respectively:

$$\mathbf{y}(t_{k+1}) = \mathbf{y}(t_k) + A(\mathbf{y}(t_k)) \mathbf{y}(t_k) h + \frac{1}{2} (A^2(\mathbf{y}(t_k)) \cdot \mathbf{y}(t_k)) h^2 + \frac{1}{2} \frac{d}{dt} [A(\mathbf{y}(t_k)) \cdot \mathbf{y}(t_k)] h^2 + O(h^3). \quad (35)$$

Taking into account that

$$\frac{d}{dt} A(\mathbf{y}(t)) = \sum_{j=1}^n \frac{\partial}{\partial y_j} A(\mathbf{y}(t)) \frac{d}{dt} y_j(t), \quad (36)$$

where y_j is the j -th component of $\mathbf{y}$, equation (35) becomes:

$$\begin{aligned} \mathbf{y}(t_{k+1}) &= \mathbf{y}(t_k) + A(\mathbf{y}_k) \mathbf{y}(t_k) h + \frac{1}{2} (A^2(\mathbf{y}(t_k)) \cdot \mathbf{y}(t_k)) h^2 \\ &+ \frac{1}{2} \left(\sum_{j=1}^n \frac{\partial}{\partial y_j} A(\mathbf{y}(t_k)) \frac{d}{dt} y_j(t_k) \right) \cdot \mathbf{y}(t_k) h^2 + o(h^3). \end{aligned} \quad (37)$$

On the k -th interval, equation (4) takes the formula

$$\mathbf{y}_N(t_{k+1}) = \mathbf{y}_N(t_k) + A(\mathbf{y}_{N_k}^*) \cdot \mathbf{y}_N(t_k) h + \frac{1}{2} A^2(\mathbf{y}_{N_k}^*) \cdot \mathbf{y}_N(t_k) h^2 + o(h^3). \quad (38)$$

Then, we may proceed to expand into Taylor series the matrix $A(\mathbf{y}_{N_k}^*)$ on a neighborhood of $\mathbf{y}_{N_k}$. Taking into account (5) and after some simple calculations, we obtain:

$$A(\mathbf{y}_{N_k}^*) = A(\mathbf{y}_N(t_k)) + \sum_{j=1}^n \frac{\partial}{\partial y_j} A(\mathbf{y}_N(t_k)) (y_{N,j}(t_k + h/2) - y_{N,j}(t_k)), \quad (39)$$

where $y_{N,j}(t)$ is the j -th component of $\mathbf{y}_N(t)$. A first order expansion on the last factor on the right hand side of (39) gives

$$y_{N,j}(t_k + h/2) - y_{N,j}(t_k) = \frac{1}{2} \frac{d}{dt} y_{N,j}(t_k) h + o(h^2), \quad (40)$$

so that using (40) in (39), we have

$$A(\mathbf{y}_{N_k}^*) = A(\mathbf{y}_N(t_k)) + \frac{h}{2} \sum_{j=1}^n \frac{\partial}{\partial y_j} A(\mathbf{y}_N(t_k)) \frac{d}{dt} y_{N,j}(t_k). \quad (41)$$

Then, we replace (41) into (38), and performing some simple manipulations, taking into account that up to second order in h :

$$\frac{1}{2} A^2(\mathbf{y}_{N_k}^*) \cdot \mathbf{y}_N(t_k) h^2 \approx \frac{1}{2} A^2(\mathbf{y}(t_k)) \cdot \mathbf{y}_N(t_k) h^2, \quad (42)$$

we finally obtain that

$$\begin{aligned} \mathbf{y}(t_{k+1}) &= \mathbf{y}_N(t_k) + A(\mathbf{y}_N(t_k)) \cdot \mathbf{y}_N(t_k) h + \frac{h^2}{2} \sum_{j=1}^n \frac{\partial}{\partial y_j} A(\mathbf{y}_N(t_k)) \frac{d}{dt} y_{N,j}(t_k) \\ &+ \frac{1}{2} A^2(\mathbf{y}_N(t_k)) \cdot \mathbf{y}_N(t_k) h^2 + o(h^3). \end{aligned} \quad (43)$$

The advantage that (43) offers with respect to (37) is that in (43) the terms up to second order in h are correctly shown.

3.1.1 Going beyond second order

In (3.1), we have found the solution up to second order in h . Would we obtain a better precision for third or higher order keeping at the same time the simplicity of the method? First of all our construction is based on equation (3), which is not longer valid if we require an approximation of order higher than two. To fix ideas, let us take $\mathbf{y}(t) = (y_1(t), y_2(t))$ bidimensional for simplicity, a choice which does not affect to our argument. Let us expand $A(\mathbf{y}(t))$ around $\mathbf{y}^*(t) = (y_1^*, y_2^*)$ (where we have omitted the sub-indices N, k for simplicity) and take one more term in the Taylor span. The result is

$$A(\mathbf{y}(t)) = A(\mathbf{y}^*(t)) + \frac{\partial}{\partial y_1} A(\mathbf{y}^*(t))(y_1(t) - y_1^*) + \frac{\partial}{\partial y_2} A(\mathbf{y}^*(t))(y_2(t) - y_2^*). \quad (44)$$

Then, instead of the approximation (3), we have the following, where again we have suppressed the indices N and k :

$$\mathbf{y}'(t) = \left[A(\mathbf{y}^*(t)) + \frac{\partial}{\partial y_1} A(\mathbf{y}^*(t))(y_1(t) - y_1^*) + \frac{\partial}{\partial y_2} A(\mathbf{y}^*(t))(y_2(t) - y_2^*) \right] \cdot \mathbf{y}(t). \quad (45)$$

The solution to be determined is just $\mathbf{y}(t) = (y_1(t), y_2(t))$, which now is a part of the Ansatz (45). One possibility to solve this contradiction is to proceed with the following span on each of the I_k intervals:

$$y_1(t) = y_1(t_k) + \frac{\partial y_1}{\partial t}(t_k)(t - t_k) + \dots, \quad (46)$$

and same for $y_2(t)$. If we use these manipulations in (3), we obtain an equation of the type:

$$\mathbf{y}'_{N,k}(t) = G_{N,k}(t) \cdot \mathbf{y}_{N,k}(t), \quad (47)$$

with

$$G_{N,k}(t) = A(\mathbf{y}_{N,k}^*(t)) + \frac{\partial}{\partial y_1} A(\mathbf{y}_{N,k}^*(t))(y_1(t) - y_1^*) + \frac{\partial}{\partial y_2} A(\mathbf{y}_{N,k}^*(t))(y_2(t) - y_2^*) + \dots \quad (48)$$

Obviously, system (48) is non-autonomic. We have to use a new approximation of $G_{N,k}(t)$ by a constant on the interval I_k and re-start again.

As we see, advancing to just one higher order of accuracy destroys the simplicity of the method which is one of its more interesting added values. Therefore, we cannot consider going to higher orders an advantage. It may slightly improve the precision at the price of destroying the efficiency and simplicity of the method.

3.2 Some examples

• The van der Pol equation

The van der Pol equation

$$y''(t) + \mu(y^2(t) - 1)y'(t) + y(t) = 0 \quad (49)$$

is a particular case of the Liénard equation,

$$y''(t) + f(y) y'(t) + g(y) = 0. \quad (50)$$

which will be discussed in the Appendix. In the van der Pol equation, we have obviously that $f(y) = \mu(y^2(t) - 1)$ and $g(y) \equiv y$. This equation can be easily written in the matrix form (2) by writing $y_1(t) \equiv y(t)$ and $y_2(t) \equiv y'(t)$, as

$$\begin{pmatrix} y_1'(t) \\ y_2'(t) \end{pmatrix} = \begin{pmatrix} 0 & 1 \\ -1 & \mu(1 - y_1^2) \end{pmatrix} \begin{pmatrix} y_1(t) \\ y_2(t) \end{pmatrix}. \quad (51)$$

Our goal is to compare the precision of our method with a reference solution. No explicit solutions to the van der Pol equation (49) are known, so that we use the Runge-Kutta solution of eight order, $y_{rk}(t)$, as reference solution (alternatively, one may consider a Taylor solution of eighth order, which has a comparable precision). We compare the precision of our method with the precision of the solutions as obtained by a third or fourth order Taylor method. As a measure of the error, we use

$$e_h := \frac{1}{N} \sum_{j=0}^{N-1} (y_{rk}(t_j) - y(t_j))^2, \quad (52)$$

where $y(t)$ is the solution obtained using our method or any other, like the Taylor method. Then, we need to use given values of the parameters and initial conditions. In Table 1 below, we compare the errors produced by our method as compared with the error given with the use of third and fourth order Taylor method for different values of the interval width h for $T = 20$, $\mu = 1/2$ and the initial conditions $y(0) = 0$ and $y'(0) = 2$.

h	Matrix	Taylor 3 rd	Taylor 4 rd
10^{-4}	$2.63 \cdot 10^{-11}$	$8.60 \cdot 10^{-7}$	$6.40 \cdot 10^{-7}$
10^{-3}	$2.08 \cdot 10^{-11}$	$8.30 \cdot 10^{-9}$	$6.45 \cdot 10^{-9}$
10^{-2}	$1.46 \cdot 10^{-8}$	$9.11 \cdot 10^{-9}$	$6.60 \cdot 10^{-11}$
10^{-1}	$2.04 \cdot 10^{-5}$	$1.17 \cdot 10^{-4}$	$1.67 \cdot 10^{-7}$
$2 \cdot 10^{-1}$	$2.30 \cdot 10^{-4}$	$2.23 \cdot 10^{-3}$	$7.94 \cdot 10^{-6}$
$5 \cdot 10^{-1}$	$8.10 \cdot 10^{-3}$	$1.49 \cdot 10^{-1}$	$2.75 \cdot 10^{-3}$

TABLE 1.- Values for the error e_h for our matrix method and the Taylor method of orders three and fourth for distinct values of h for the van der Pol equation.

It is clear that our matrix method has a precision in between those of the third and fourth order Taylor method. These results are just an example of the results obtained in multiple numerical experiments, we have performed showing essentially the same result. However, Table 1 as well as other numerical experiments show an important tendency: the lower h

the better our precision as compared with the precision given by the Taylor method. The reason is clear, the smaller h the bigger the number of operations needed to obtain the approximate solution. Our matricidal method requires less operations than the Taylor method, so that our precision gets better as h gets smaller.

In Appendix B, we give our source code when our method is to be applied in the present case.

- **The Duffing equation**

This is another second order non linear equation, which has the following form:

$$y''(t) + y(t) + y^3(t) = 0, \quad (53)$$

where we have omitted the term in the first derivative of the indeterminate, $y'(t)$. For the Duffing equation, there are explicit solutions in terms of the Jacobi elliptic functions. For instance, being given the initial conditions $y(0) = 1$ and $y'(0) = 0$, we have the following solution:

$$y_e(t) = -i\sqrt{1+k} \operatorname{sn}(u; m), \quad (54)$$

where $\operatorname{sn}(u; m)$ denotes the elliptic sine. The arguments in (49) denote the following:

$$u = \frac{1}{\sqrt{2}} (t^2(1-k) + 2t(c_2 - kc_1) + (1-k)c_2^2),$$

$$m = \frac{1+k}{1-k}, \quad k = \sqrt{1+2c_1}. \quad (55)$$

The value of the constants included in (55) are $c_1 = 1.5$ and $c_2 = 1.1920055$. The Duffing equation may be written in matrix form, if we write again $y_1(t) \equiv y(t)$ and $y_2(t) \equiv y'(t)$, as

$$\begin{pmatrix} y_1'(t) \\ y_2'(t) \end{pmatrix} = \begin{pmatrix} 0 & 1 \\ -(1+y_1^2) & 0 \end{pmatrix} \begin{pmatrix} y_1(t) \\ y_2(t) \end{pmatrix}. \quad (56)$$

We define the error e_h as in (52), where we replace $y_{rk}(t)$ by the exact solution, $y_e(t)$, which does exist in the present case. Also $y(t)$ is the solution for which we want to compare its precision with the exact solution, in our case the solutions obtained by our matrix method as well as the third or fourth order Taylor solutions. The errors produced by each method are given on Table 2 below.

h	Matrix	Taylor 3 rd	Taylor 4 rd
10^{-4}	$1.69 \cdot 10^{-9}$	$5.60 \cdot 10^{-5}$	$4.13 \cdot 10^{-5}$
10^{-3}	$1.03 \cdot 10^{-10}$	$5.06 \cdot 10^{-7}$	$4.24 \cdot 10^{-7}$
10^{-2}	$9.61 \cdot 10^{-9}$	$2.47 \cdot 10^{-8}$	$4.24 \cdot 10^{-9}$
10^{-1}	$1.16 \cdot 10^{-5}$	$4.92 \cdot 10^{-4}$	$4.12 \cdot 10^{-8}$
$2 \cdot 10^{-1}$	$1.18 \cdot 10^{-4}$	$6.78 \cdot 10^{-3}$	$7.56 \cdot 10^{-6}$
$5 \cdot 10^{-1}$	$82.00 \cdot 10^{-3}$	$1.04 \cdot 10^{-1}$	$7.73 \cdot 10^{-3}$

TABLE 2.- Values for the error e_h for our matrix method and the Taylor method of orders three and fourth for distinct values of h for the Duffing equation.

We see that the results are quite similar to those obtained with the van der Pol equation. Similarly, we have made some numerical experiments that confirm these results.

• The Lorenz equation

The Lorenz equation, which is a model for the study of chaotic systems has been introduced in the study of atmospheric behaviour [9]. This equations arises in many problems of physics, where chaoticity is present [17–20]. The Lorenz equation is usually written in matrix form and is three-dimensional:

$$\begin{pmatrix} y_1'(t) \\ y_2'(t) \\ y_3'(t) \end{pmatrix} = \begin{pmatrix} -a & a & 0 \\ b - y_3(t) & -1 & 0 \\ y_2(t) & 0 & -c \end{pmatrix} \begin{pmatrix} y_1(t) \\ y_2(t) \\ y_3(t) \end{pmatrix}, \quad (57)$$

a , b and c being positive constants.

This system is very sensitive to the particular choice of the parameters and of initial values, as may numerical experiments show. Based in these experiments, we have choosed the following values for the parameters, $a = 10$, $b = 9.996$ and $c = 8/3$, the fixed point $(y_1, y_2, y_3) = (4.88808, 4.88808, 8.996)$ is an attractor. We proceed as in the previous case and compare solutions of the errors (52) resulting of the use of the Taylor method of order two and this Matrix method using as reference the solution obtained with the Ruge Kutta method of eighth order. We use $T = 10$ and the solution with initial values $(y_1(0), y_2(0), y_3(0)) = (1, 0, 1)$. We have obtained a table (Table 3) of errors given in terms of the interval width h as

h	Matrix	Taylor second order	Taylor third order	Taylor fourth order
10^{-2}	$5.2 \cdot 10^{-6}$	$6.73 \cdot 10^{-5}$	$3.6 \cdot 10^{-8}$	$3.1 \cdot 10^{-10}$
10^{-1}	$5.3 \cdot 10^{-2}$	$7.1 \cdot 10^{-1}$	$1.8 \cdot 10^{-1}$	$5.2 \cdot 10^{-2}$
$2 \cdot 10^{-1}$	$3.7 \cdot 10^{-1}$	error	error	error

TABLE 3.- Values of the precision e_h in terms of h for $T = 10$, $y_1(0) = 1$, $y_2 = 0$ and $y_3(0) = 1$.

The word “error” written in two entries in Table 2 means that the error that appears if the interval width is of the order of 0.2, when we use the Taylor method, is incontrollable. We also see that the precision obtained with the matrix method clearly improves the precision by the Taylor method the larger the length of the subintervals.

- **Neutral damping equation.**

This equation has been discussed in the literature [10, 16] and is

$$x''(t) + \varepsilon (x'(t))^2 + x(t) = 0, \quad (58)$$

where the tilde means derivative with respect to the variable t and ε is a real parameter. As in previous cases, let us define $y(t) := x'(t)$, so that (58) can be written in the following form:

$$(-x + \varepsilon y^2) dx - y dy = 0. \quad (59)$$

This equation is integrable with integrating factor:

$$\mu(x, y) \equiv e^{2\varepsilon x}. \quad (60)$$

It is readily shown that equation (60) admits the following first integral:

$$f(x, y) = \left[\frac{1}{2} y^2 + \frac{1}{4\varepsilon^2} (2\varepsilon x - 1) \right] e^{2\varepsilon x} \quad (61)$$

If we write (58) in the standard matrix form as

$$\begin{pmatrix} x' \\ y' \end{pmatrix} = \begin{pmatrix} 0 & 1 \\ -1 & -\varepsilon y \end{pmatrix} \begin{pmatrix} x \\ y \end{pmatrix}, \quad (62)$$

we readily observe that the only fixed point is the origin. It is also the origin the point at which the minimum of the first integral (61) lies. All orbits are periodic around the origin.

In Figure 1, we have depicted the periodic solutions in phase space. Note that, although equation (58) may look like dissipative, it is not as Figure 1 manifests.

Now, let us repeat the comparison between errors given by our matrix method with the errors given by the Taylor method at some low orders. As always, we use (47) as the definition of the error and the numerical Runge-Kutta solution of eighth order as the reference solution. We obtained the following results, which appear in Table 4, where h is, as always, the distance between two consecutive nodes:

h	Matrix	Taylor 3 rd	Taylor 4 rd
10^{-4}	$2.2 \cdot 10^{-14}$	$4.2 \cdot 10^{-14}$	$3.2 \cdot 10^{-11}$
10^{-3}	$1.9 \cdot 10^{-14}$	$4.1 \cdot 10^{-11}$	$3.1 \cdot 10^{-11}$
10^{-2}	$9.0 \cdot 10^{-14}$	$4.6 \cdot 10^{-13}$	$4.0 \cdot 10^{-12}$
10^{-1}	$1.6 \cdot 10^{-9}$	$2.7 \cdot 10^{-7}$	$8.2 \cdot 10^{-11}$
$2 \cdot 10^{-1}$	$2.6 \cdot 10^{-8}$	$8.6 \cdot 10^{-6}$	$1.0 \cdot 10^{-9}$
$5 \cdot 10^{-1}$	$1.1 \cdot 10^{-6}$	$7.5 \cdot 10^{-4}$	$6.1 \cdot 10^{-6}$

TABLE 4. Values of the error in terms of h for $\varepsilon = 0.1$, $T = 2\pi$ (see (47)) and the initial values $x(0) = 1$, $x'(0) = 0$.

We observe that the precision of the matrix method is much higher than the precision of the third and forth order Taylor method for a distance between nodes $h \leq 0.01$. Moreover, we have to underline that our matrix one step method is much simpler to programming than the Taylor method as the reader can easily convince him/herself using any of these examples.

There is another method based in the theory of perturbations in order to obtain approximate solutions to (53), which receives the name of Lindstedt-Poincaré. It is described in [16]. It consists in a series in terms of ε of the form:

$$x(t) = x_0(t) + \varepsilon x_1(t) + \varepsilon^2 x_2(t) + \dots \quad (63)$$

Coefficients $x_i(t)$ may be obtain iteratively, once we have fixed the initial values. For instance, for $x(0) = 1$ and $x'(0) = 0$, we obtain:

$$\begin{aligned} x_0(t) &= \cos \omega t, & x_1(t) &= \frac{1}{6} (-3 + 4 \cos \omega t - \cos 2\omega t), \\ x_2(t) &= \frac{1}{3} \left(-2 + \frac{61}{24} \cos \omega t - \frac{2}{3} \cos 2\omega t + \frac{1}{8} \cos 3\omega t \right), \\ \omega &= 1 - \frac{1}{6} \varepsilon + o(\varepsilon^3). \end{aligned} \quad (64)$$

We may also evaluate the error for this perturbative method, which is independent of any division of the interval, in which the solution is considered, into subintervals. This error is $1.1 \cdot 10^{-4}$, which is obviously higher to those obtained using any of the numerical method considered.

We finish the present example by proposing another approach to an approximate solution. By either method, matrix or Taylor, we obtain on each of the nodes $\{t_k\}$ a value, x_k , of the approximate solution. Let us interpolate each interval by means of cubic splines, so

that we obtain a segmentary approximation by cubic polynomials ². Let us assume that the cubic interpolating solution is $s(t)$, then the error of this solution on an interval T , with respect to the exact solution is given by (recall that cubic splines admit first and second continuous derivatives at the nodes)

$$e = \frac{1}{T} \int_0^T (s''(t) + \varepsilon [s'(t)]^2 + s(t))^2 dt. \quad (65)$$

The form of the error for the Taylor method is also given by (60). The resulting errors appear in the following table (Table 5):

h	Spline	Taylor 2 rd	Taylor 3 rd	Taylor 4 rd
10^{-4}	$7.2 \cdot 10^{-16}$	$4.7 \cdot 10^{-16}$	$6.7 \cdot 10^{-16}$	$6.5 \cdot 10^{-16}$
10^{-3}	$1.8 \cdot 10^{-14}$	$1.8 \cdot 10^{-14}$	$1.8 \cdot 10^{-14}$	$1.8 \cdot 10^{-14}$
10^{-2}	$1.9 \cdot 10^{-10}$	$1.9 \cdot 10^{-11}$	$1.1 \cdot 10^{-10}$	$1.1 \cdot 10^{-10}$
10^{-1}	$4.6 \cdot 10^{-6}$	$6.4 \cdot 10^{-6}$	$4.6 \cdot 10^{-6}$	$4.6 \cdot 10^{-6}$
$2 \cdot 10^{-1}$	$6.1 \cdot 10^{-5}$	$9.8 \cdot 10^{-5}$	$6.0 \cdot 10^{-5}$	$6.0 \cdot 10^{-5}$
$5 \cdot 10^{-1}$	$4.8 \cdot 10^{-3}$	$1.2 \cdot 10^{-3}$	$4.6 \cdot 10^{-3}$	$4.9 \cdot 10^{-3}$

TABLE 5.- Comparison between errors by the cubic spline and Taylor methods.

Finally, for the perturbative method the error obtained is $7.4 \cdot 10^{-5}$. Concerning the conservation of the constant of motion, we measure its dispersion by means of the following parameter, e_f , defined as

$$e_f := \frac{1}{T} \int_0^T (f(x_0, y_0) - f(x, y))^2 dt, \quad (66)$$

where $f(x, y)$ should be calculated using different approximations such as Taylor, matrix and the analytic approximate solution as obtained by the perturbative Lindstedt-Poincaré method mentioned earlier. The point (x_0, y_0) gives the chosen initial conditions. In the latter case, we have obtained $e_f = 1.1 \cdot 10^{-4}$, in all others, we always got $e_f < 10^{-8}$.

• An epidemic equation

A model for an epidemia has been proposed as early as in 1927 [23–25]. If $x(t)$, $y(t)$ and $z(t)$ are the number, at certain time t , of healthy, sick and dead persons, respectively, in some society, the model assumes that these functions satisfy the following non-linear system:

²In [22], we have already proposed segmentary cubic solutions and studied their properties, although in [22] they were not necessarily cubic splines.

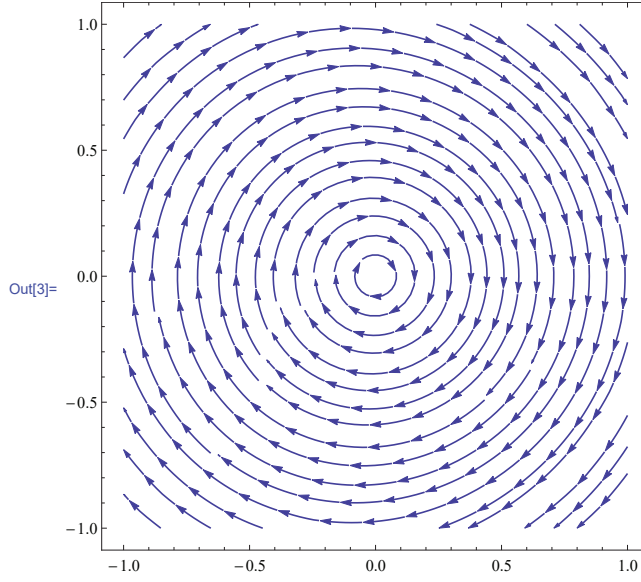


Figure 1: Periodic orbits around the origin for the neutral damping equation. The horizontal coordinate represents the values of x , while vertical coordinate gives the values of $y = x'$. Note that these periodic orbits are represented in phase space.

$$\begin{aligned}
 \dot{x}(t) &= -ax(t)y(t), \\
 \dot{y}(t) &= ax(t)y(t) - by(t), \\
 \dot{z}(t) &= by(t),
 \end{aligned} \tag{67}$$

where a and b are positive constants and the dot means derivation with respect to time t . The model assumes infection of healthy persons from sick persons. The latter died after some time. Note that, in the studied time interval, the sole cause of population dynamics is the epidemic. For obvious reasons, we consider only positive solutions. The fixed points for (67) have the form $(\alpha, 0, \beta)$ with $\alpha, \beta > 0$.

Let us consider the following vector field, also called the *flux* of (62),

$$F(t) := (-ax(t)y(t), ax(t)y(t) - by(t), by(t)). \tag{68}$$

Note that $F_x < 0$. For $x(t) < b/a$, we have $F_y < 0$, while $x > b/a$, it results that $F_y > 0$. Thus, the fixed point is inestable if $\alpha > b/a$. This is a necessary condition for the existence of the epidemic. If $\alpha < b/a$, then the fixed points are asymptotically stable as depicted in Figure 2, where we have chosen $a = 0.0005$ and $b = 0.1$. Different curves in Figure 2 represent different initial conditions.

Let us write (67) in matricial form as

$$\dot{X}(t) = A(x, y, z) \cdot X(t), \tag{69}$$

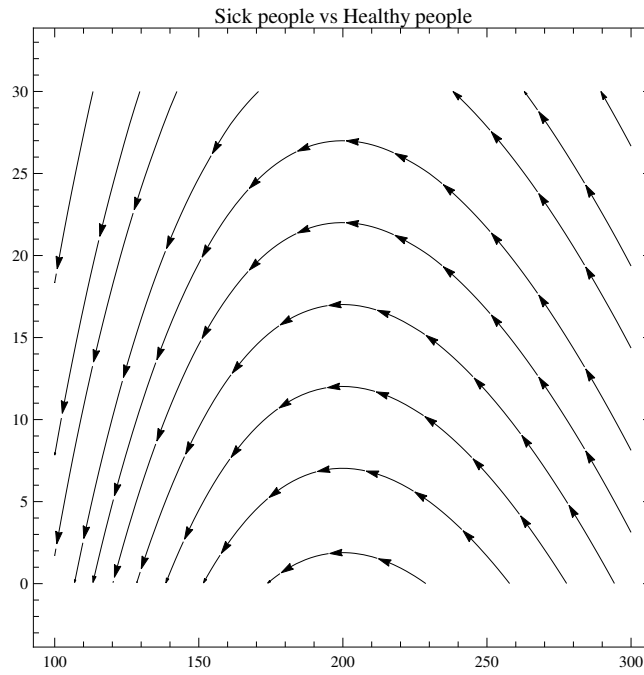


Figure 2: Asymptotic stability of fixed points with $a = 0.0005$ and $b = 0.1$ in (67). The horizontal line represent the scaled number of healthy people, while the vertical figures give the scaled number of sick people. The arrow means the direction of time. Different curves are obtained using different initial conditions. Observe the presence of a maximum at a point which is independent on the initial conditions.

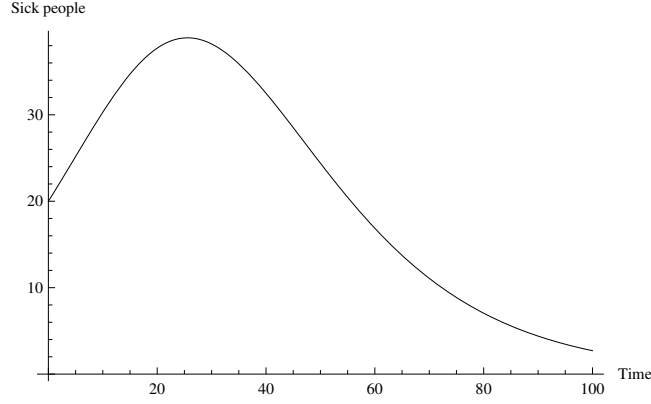


Figure 3: Number of infected people in relation with time. Observe the existence of a maximum. After the maximum the curve decreases steeply.

where,

$$X(t) := \begin{pmatrix} x(t) \\ y(t) \\ z(t) \end{pmatrix}, \quad A(x, y, z) := \begin{pmatrix} 0 & -ax & 0 \\ 0 & ax - b & 0 \\ 0 & b & 0 \end{pmatrix}. \quad (70)$$

In Figure 3, we obtain the curve giving the total number of infected people with time. After a maximum, the number of sick people decays quickly. We have used the values for the parameters $a = 0.0005$ and $b = 0.1$ and the initial values $x(0) = 300 > b/a$, $y(0) = 20$, $z(0) = 0$.

The form of the error for the method is obtained using (65) again. Next in Table 6, we compare the errors for given values of h of our Matrix Method as compared to second and third order Taylor.

h	Matrix Method	Taylor 2 rd	Taylor 3 rd
0.01	$2.2 \cdot 10^{-11}$	$4.0 \cdot 10^{-11}$	$1.9 \cdot 10^{-11}$
0.1	$2.0 \cdot 10^{-8}$	$7.3 \cdot 10^{-8}$	$2.0 \cdot 10^{-11}$
1.0	$2.1 \cdot 10^{-4}$	$7.2 \cdot 10^{-4}$	$1.7 \cdot 10^{-7}$
2.0	$3.6 \cdot 10^{-3}$	$1.1 \cdot 10^{-2}$	$1.2 \cdot 10^{-5}$

TABLE 6.- Comparison between the errors for the Matrix Method and the Taylor method in the case of the epidemic model.

We see that the level of error is similar, although our method keeps the advantage of needing much less arithmetics than any Taylor method.

- **Lotka-Volterra equation.**

Models in population dynamics, as for instance the predator-prey competition, were independently developed by the american biologist A.J. Lotka [11] and the Italian mathematician V. Volterra [12], see modern references for the Lotka-Volterra equation in [4, 26–28]. The most general form of the Lotka-Volterra equation has the form

$$\begin{aligned}\dot{x}_1 &= x_1(\varepsilon_1 - a_{11} x_1 - a_{12} x_2), \\ \dot{x}_2 &= x_2(\varepsilon_2 - a_{12} x_1 - a_{22} x_2),\end{aligned}\tag{71}$$

where x_1 and x_2 are functions of time t and ε_i , a_{ij} , $i, j = 1, 2$ are constants. Following [26], we consider here a simpler version, yet non-linear, of (71), which is

$$\begin{aligned}\dot{x} &= a x - b x y, \\ \dot{y} &= d x y - c y,\end{aligned}\tag{72}$$

and the initial conditions $x(t_0) = x_0$, $y(t_0) = y_0$.

We may consider the solutions, $x = x(t)$, $y = y(t)$ of (72) as the equations determining a parametric curve on the $x - y$ plane. Then, by elimination of t and integration, we obtain:

$$C(x, y) = x^c y^a \exp\{-(b + d)x\}.\tag{73}$$

Note that $C(x, y)$ is a constant on each curve solution (constant of motion) and its value over each solution is determined via the initial conditions. Equations (71) have two fixed points, which are $(0, 0)$ and $(c/d, a/b)$. For $x > 0$ and $y > 0$ all solutions are periodic.

To apply the method proposed in the present article to this system, let us write it in matrix form as

$$\dot{X} = A(x, y) X,\tag{74}$$

where,

$$X(t) \equiv \begin{pmatrix} x(t) \\ y(t) \end{pmatrix}, \quad A \equiv \begin{pmatrix} a - by & 0 \\ 0 & dx - c \end{pmatrix}.\tag{75}$$

In order to estimate the precision of the method, we need to choose values for the initial values as well as for the parameters a , b , c and d . We have use several choices and obtain in all of them similar values. to show a table comparing the precision of our method with some others, let us choose as the values of the parameters, $a = 1.2$, $b = 0.6$, $c = 0.8$ and $d = 0.3$. As the initial values, let us choose, $x_0 = y_0 = 3$. In addition, we have to take an integration time, in our case we took $T = 20$, which approximately accounts for three periods.

In order to determine the error produced by our matrix method is determined by the formula (65) above. We compare this error with those of second and third order Taylor method as compared to the numerical solution obtained by a forth order Runge-Kutta. These errors are shown in Table 7.

h	Matrix Method	Taylor 2 rd	Taylor 3 rd
0.01	$4.1 \cdot 10^{-8}$	$3.5 \cdot 10^{-8}$	$8.2 \cdot 10^{-13}$
0.1	$3.5 \cdot 10^{-8}$	$5.4 \cdot 10^{-8}$	$2.3 \cdot 10^{-9}$
0.2	$2.6 \cdot 10^{-3}$	$5.4 \cdot 10^{-3}$	$1.0 \cdot 10^{-5}$

TABLE 7.- Comparison between the errors for the Matrix Method and the Taylor method in the case of the Lotka-Volterra equation.

We see that the precision of our method is equivalent to those of a second order Taylor, with much less arithmetic operations.

- **On the possibility of extending the method to PDE: The Burger's equation.**

Can we extend this precedent discussion to partial differential equations admitting separation of variables? One possible example would have been the convection diffusion equation in one spatial dimension [29]:

$$\frac{\partial}{\partial t} u(x, t) = \frac{\partial}{\partial x} \left[D(u) \frac{\partial}{\partial x} u(x, t) \right] + Q(x, u) \frac{\partial}{\partial x} u(x, t) + P(x, u). \quad (76)$$

Separation of variables for (76) is discussed in [29]. In general, our method is not applicable here since the resulting equations after separation of variables are not of the form (1). Nevertheless, another point of view is possible. Assume we want to obtain approximate solutions of an equation of the type (76) on the interval $[0, X]$ under the conditions $u(0, t) = 0$, $u(X, t) = 0$, $u(x, 0) = h(x)$, $h(x)$ being a given smooth function and $u(0, 0) = u(X, 0) = 0$. On the interval $[0, X]$, we define a uniform partition of width $h := X/n$ and nodes $x_k = kh$, $k = 0, 1, \dots, n$.

One may propose one discretization of the solution of the form $u_k(t) := u(x_k, t)$, $k = 0, 1, \dots, n$. Then, the second derivative in (76) could be approximated using finite differences [30]:

$$\frac{\partial^2}{\partial x^2} u(x_k, t) = \frac{u(x_{k-1}, t) - 2u(x_k, t) + u(x_{k+1}, t))}{h^2}, \quad (77)$$

while for the first spatial derivative, we have

$$\frac{\partial}{\partial x} u(x_k, t) = \frac{u(x_{k+1}, t) - u(x_{k-1}, t))}{2h}. \quad (78)$$

with $k = 1, 2, \dots, n-1$, so that equation (76) takes the form

$$\frac{d}{dt} U(t) = F(U(t)), \quad (79)$$

where F is a square matrix of order $n-1$ and $U(t) = (u_1(t), u_2(t), \dots, u_{n-1}(t))$ with $u_k(t) := u(x_k, t)$, $k = 1, 2, \dots, n-1$ and initial values $U(0) = (u(x_1, 0), u(x_2, 0), \dots, u(x_{n-1}, 0))$ and initial value $u(x, 0) = h(x)$. If it were possible to write equation (79) in the form:

$$\frac{d}{dt} U(t) = A(U(t)) \cdot U(t), \quad (80)$$

then, we would be able to apply our method to find an approximate solution of (80) on a given finite interval. This property is not fulfilled by the general convection diffusion equation (76). However, it is satisfied by a particular choice of this type of parabolic equations: *the non-linear Burger's diffusion equation*, which is [31, 32]:

$$\frac{\partial}{\partial t} u(x, t) = \frac{\partial^2}{\partial x^2} u(x, t) - u(x, t) \frac{\partial}{\partial x} u(x, t). \quad (81)$$

We choose $[0, 1]$ as the integration interval for the coordinate variable, so that $X = 1$ as above. For the time variable, we use $t \in [0, 1]$. This choice is made just for simplicity and as an example to implement our numerical calculations. We use the finite difference method, where the second derivative and first derivatives are replaced as in (77) and (78), respectively.

Using (77) and (78) in (81), we have for each of the nodes the following recurrence relation:

$$\frac{d}{dt} u(x_k, t) = \frac{1}{h^2} \left(\left(1 - \frac{h}{2} u(x_k, t) \right) u(x_{k-1}, t) - 2u(x_k, t) + \left(1 + \frac{h}{2} u(x_k, t) \right) u(x_{k+1}, t) \right), \quad (82)$$

for $k = 1, 2, \dots, n-1$. When written expressions (82) in matrix form, we obtain a matrix equation of the form (6) and, then, suitable for applying our method.

In our numerical realization, we have used a small number of nodes, to begin with, say $n = 5$. Then, we integrate on the time variable, on the interval $t \in [0, 1]$, using $h_t := 1/m$, where m is a given integer, as the distance between time nodes. Then, we compare our solution with the solution of (81) given by the sentence `NDSolve` provided by the Mathematica software, solution that we denote as $v(x, t)$.

Then, we can estimate the error of our solution as compared with $v(x, t)$. This is given by

$$\text{error} := \frac{1}{m} \sum_{j=0}^m \sum_{k=1}^{n-1} (u(x_k, t_j) - v(x_k, t_j))^2, \quad (83)$$

where $t_j := jh_t$ for all $j = 0, 1, 2, \dots, m$. To estimate the error, we have to give values to m . For $m = 10$, the error is $1.07 \cdot 10^{-4}$. A similar result can be obtained with $m = 20$ or even higher, so that $m = 10$ gives already a reasonable approximation. The solution

for $t = 1$ is nearly zero, which one may have expected taking into account that equation (80) describes a dissipative model. Thus, our approximate solution may be considered satisfactory also in this case.

A few more words on the comparison of our solution and the solution using the Euler method, performed through NDSolve. First of all, we use the explicit Euler method, for which the local error is $o(h_t^2)$, through the option “Method $\rightarrow$ “ExplicitEuler”, “StartingStepSize” $\rightarrow 1/100$ ”, since for $h_t > 0.01$ the result is unstable. Once we have done the spatial discretization, we integrate (82) with respect to time. The instability often appears when one uses an explicit method of spatial and time discretization for parabolic PDE [33]. This instability comes after the errors due to the arithmetic calculations and are amplified after time integration. Then, let us consider $n = 5$, where n is the number of spatial nodes, $0 \leq t \leq 1$ and $m = 100$, so that the time integration interval becomes $h_t = 0.01$. Then, the error (83) is $3.40 \cdot 10^{-4}$.

We may improve time integration using Euler mid-point integration. In this case, one uses the option “Method $\rightarrow$ “ExplicitMidpoint”, “StartingStepSize” $\rightarrow 1/10$ ”. Choosing $m = 10$, we obtain an error of $1.76 \cdot 10^{-4}$, so that $h_t = 0.1$. This error is of the order given by our method. Just recalling that the local error given by the mid-point Euler method is of the order of $o(h_t^3)$.

In addition, we may compare the precision of our method and both Euler methods mentioned above by the errors obtained using the third and forth order Adams method (see Chapter III in [34]), which for $h_t = 0.1$ are respectively given by $1.76 \cdot 10^{-4}$ and $1.74 \cdot 10^{-4}$. This error has always been obtained using (83), where $u(x, t)$ is our solution and $v(x, t)$ is the solution given by either Euler, Euler mid-point or Adams methods. Nevertheless, the Adams method is a multi-step method while ours is one step method, which means that ours is more easily programmable.

4 Concluding remarks

In the present paper, we have generalized a one step integration method, which has been developed in the seventies of last century by Demidovich and some other authors [1–3]. While Demidovich and others restrict themselves to the search for approximate solutions of linear systems of first order differential equations, albeit with variable coefficients, we propose a way to extend the ideas of the mentioned authors to non-linear systems. The solutions we have obtained have a similar degree of precision than those proposed in [1–3]. In our method, we obtain a sequence of approximate solutions on a finite interval of ordinary differential equations, and we have proven that this sequence converges uniformly to the exact solution. In order to obtain each approximate solution, we have divided the integration interval into subintervals of length h . The sequence of approximate solutions can be obtained after successive refinements of h . For each approximate solution, characterized by a value of h , we have determined the local error up to $o(h^3)$.

It is certainly true, that our one step matrix method does not improve the precision obtained with the fourth order Runge-Kutta method or other equivalent. However, the great advantage of the proposed method with any other is that its algorithm of construction, through the

exponential matrix as described in Section 2, is much simpler than their competitors. Simplicity that is inherited from its antecesor the Demidovich method [1–3]. This paper is somehow the continuation of previous research of the authors in the same field [35, 36].

We have added some examples of the application of the method, on where we have performed numerous numerical test on the precision of the method, which lies between second and third Taylor method. We have used the software Mathematica to implement these numerical tests.

Acknowledgements

We acknowledge partial financial support to the Spanish MINECO, grant MTM2014-57129-C2-1-P, and the Junta de Castilla y León, grants VA137G18, BU229P18 and VA057U16. We are grateful to Prof L.M. Nieto (Valladolid) for some useful suggestions.

5 Appendix A: A conjecture relative to the Liénard equation.

As is well known, the Liénard equation has the following form:

$$y''(t) + f(y) y'(t) + g(y) = 0. \quad (84)$$

Let us assume that the function $g(y)$ is a product of some function, that we also call $g(y)$ for simplicity, and y , so that equation (67) takes the form:

$$y''(t) + f(y) y'(t) + g(y) y(t) = 0. \quad (85)$$

This second order equation may be easily transformed into a two dimensional first order equation ($y(t) = y(t)$, $z(t) = y'(t)$):

$$\begin{pmatrix} y' \\ z' \end{pmatrix} = \begin{pmatrix} 0 & 1 \\ -g(y) & -f(y) \end{pmatrix} \begin{pmatrix} y \\ z \end{pmatrix} = A \begin{pmatrix} y \\ z \end{pmatrix}, \quad (86)$$

where the meaning of the matrix A is obvious. Its eigenvalues are

$$\lambda_{\pm}(y) = -\frac{1}{2} \left(f(y) \pm \sqrt{f^2(y) - 4g(y)} \right). \quad (87)$$

The conjecture is the following: A sufficient condition for the solutions of (67) to be bounded for $t > 0$ is that the following three properties hold simultaneously: i.) The discriminant in (60) be negative, i.e., $f^2(y) - 4g(y) < 0$; ii.) The function $f(y)$ be non-negative, $f(y) \geq 0$ and iii.) The function $g(y)$ is positive and smaller than one, $0 < g(y) < 1$.

This conjecture is based in the form of equations (11) and (12) and we have tested it in several numerical experiments.

Two more comments in relation to the Liénard equation:

1.- The vector field associated to the equation is given by $J \equiv (z, -g(y)y - f(y)z)$, for which the divergence is $\text{div}(J) = -f(y) \leq 0$ if $f(y) \geq 0$. Therefore, if $f(y)$ is non-negative the

divergence is always negative, so that the origin $(0, 0)$ is an attractor. We conjecture that this attractor is also global.

2.- If we take $f(y) \equiv 0$, then (72) represents a Hamiltonian flow with Hamiltonian given by

$$H(y, z) = \frac{1}{2} z^2 + V(y) , \quad (88)$$

with

$$V(y) = \int_0^y u g(u) du \quad (89)$$

Under the assumption $g(y) > 0$, the derivative $V'(y) = yg(y)$ is positive for $y > 0$ and negative for $y < 0$. Thus the only critical point is $y^* = 0$, at this point $V'(0) = 0$ and $V''(0) = g(0) > 0$, so that all the orbits are closed, hence periodic.

6 Appendix B: Source code to use our method in the van der Pole equation

```

(* Example:  van der Pol equation *)
 $\mu = 1/2$ ; (* parameter*)
 $y_0 = \{0, 2\}$ ; (* initial value *)
xmax = 10.; (* integration interval [0, xmax] *)
n = 1000;    (* number of interval *)
h = xmax/n  (* integration step *)
A := {{0, 1}, {-1, + $\mu$  (1 - (Part[yk, 1])2)}} (* matrix *)
B := {{0, 1}, {-1, + $\mu$  (1 - (Part[zk+1, 1])2)}} (* auxiliary matrix*)
Do[{zk+1 = Simplify[MatrixExp[A  $\frac{h}{2}$ ].yk], yk+1 = Simplify[MatrixExp[B h].yk]}
  {k, 0, n}] (* step by step iterations*)
tabla = Table[{k h, Part[yk, 1]}, {k, 1, n}]; (* Table with the solution *)

```

References

- [1] Demidovich B. *Lectures on the mathematical Theory of Stability*, Nauka, Moscow, 1967.
- [2] Kotsis D. The approximation of the characteristic multipliers of periodic differential equations. *Alkalmaz. Mat. Lapok.* 1977; 2: 269-276.
- [3] Hsu C.S. On approximating a general linear periodic system. *J. Math. Anal. Appl.* 1974; 45: 234-251.
- [4] Farkas M. *Periodic Motions*. Springer, New York, 1994.
- [5] van der Pol B. On relaxation-oscillations, *The London, Edinburgh and Dublin Phil. Mag. & J. of Sci.* 1927; 2 (7): 978-992.
- [6] Nepomuceno E.P., Peixoto M.L.C., Martins S.A.M., Rodrigues Junior H.M., Perc M., Inconsistencies in numerical simulations of dynamical systems using interval arithmetic, *Int. J. Bifur. Chaos* 2018; 28 (04): 1850055.
- [7] Nepomuceno E.P., Guedes P.F.S., Barbosa A.M., Perc M., Repnik R., Soft computing simulations of chaotic systems, *Int. J. Bifur. Chaos* 2019; 29 (08): 1950112.
- [8] Duffing G. *Forced oscillations with variable natural frequency and their technical relevance* (German), Heft 41/42, Braunschweig: Vieweg; 1918, 1-134.
- [9] Lorenz E.N., Deterministic non-periodic flow, *J. Atmos. Sci.* 1963; 20 (2): 130-141.
- [10] Hale J., Kocak H., *Dynamics and Bifurcations*. Springer-Verlag; New York, 1991.
- [11] Lotka A.J. *Elements of Physical Biology*. Williams and Wilkins; Philadelphia, USA, 1925.
- [12] Volterra V. Variations and fluctuations of the number of individuals in animal species living together, in *Animal Ecology*, R.N. Chapman, Ed., McGraw-Hill, New York, 1931.
- [13] Beléndez A., Gimeno E., Fernández E., Méndez D.I., Alvarez M.L. Accurate approximate solutions to non-linear oscillators in which the restoring force is inversely proportional to the dependent variable. *Phys. Scr.* 2008; 77: 065004 (2008).
- [14] Egger H., Schmidt K., Shashkov V. Multistep and RungeKutta convolution quadrature methods for coupled dynamical systems, to be published in the *J. Comp. Appl. Math.*
- [15] Kalogiratou Z., Monovasilis Th., Psihoyios G., Simos T.E., Runge-Kutta type methods with special properties for the numerical integration of ordinary differential equations. *Phys. Rep.* 2014; 536: 75-146.
- [16] Mickens R. *Oscillations in Planar Dynamic Systems*. World Scientific, London, 1996.
- [17] Haken H. Analogy between higher instabilities in fluids and lasers. *Phys. Lett. A* 1975; 53: 77-78.

- [18] Gorman M., Widmann P.J., Robbins K.A., Nonlinear dynamics of a convection loop: A quantitative comparison of experiment with theory. *Physica D*. 1986; 19 (2): 255-267.
- [19] Coumo K.M., Oppenheim A.V., Circuit implementation of synchronized chaos with applications to communications. *Phys. Rev. Lett.* 1993; 71 (1): 65-68.
- [20] Mishra A., Sanghi S. A study of the asymmetric Malkus waterwheel: The biased Lorenz equations. *Chaos* 2006; 16 (1): 013114.
- [21] Bario R. Performance of the Taylor series method for ODEs/DAEs. *Appl. Math. Comp.* 2005; 163: 525-545.
- [22] Lara L.P., Gadella M. An approximation to solutions of linear ODE by cubic interpolation *Comp. Math. Appl.* 2008; 56: 1488-1495.
- [23] Kermack W.O., McKendrick A.G. A contribution to the Mathematic Theory of Epidemics. *Proceedings of the Royal Society A* 1927; 115 (772): 700-721.
- [24] S. Strogatz. *Non-linear Dynamics and Chaos*. Addison-Wesley, New York, 1994.
- [25] Torralba-Rodríguez O., Conde-Gutiérrez R.A., Hernández-Javier A.L. Modeling and prediction of COVID-19 in México applying mathematical and computational models. *Chaos, Solitons and Fractals* 2020; 138: 1009946 (2020).
- [26] Chicone C. *Ordinary Differential Equations with Applications*. Springer, New York, 1999.
- [27] Murray J.D. *Mathematical Biology I: An Introduction*. Springer, Berlin and New York, 2003.
- [28] Llibre J., Valls C., Global analytic first integrals for the real planar Lotka-Volterra system. *J. Math. Phys.* 2007; 48 (3): 033507.
- [29] Huabing Jia, Wei Xu, Xiaoshan Zhao, Zhanguo Li, Separation of variables and exact solutions to nonlinear diffusion equations with x-dependent convection and absorption. *J. Math. Anal. Appl.* 2008; 339: 982-995.
- [30] Kincaid D, Cheney W, *Numerical analysis, Mathematics of Scientific Computing*, AMS, Providence, Rhode Island, USA, Third Edition 2002.
- [31] Burger J.M., *A mathematical Model Illustrating the Theory of Turbulence*, Academic Press, New York, 1948.
- [32] Biazar J., Aminikhah H., Exact numerical solutions for non-linear Burger's equation by VIM. *Math. Comp. Mod.* 2009; 49: 1394-1400.
- [33] Carnahan B., Lutter H.A., Wilkes J.O., *Applied Numerical Methods*. Wiley, New York, 1969.
- [34] Hairer E., Nørsett S.P., Wanner G., *Solving Ordinary Differential equations I: Nonstiff Problems*. Springer, Berlin, Heidelberg, 1993.

- [35] Gadella M., Lara L.P. On the determination of approximate periodic solutions of some non-linear ODE. *Appl. Math. Comp.* 2012; 18 (10): 6038-6044.
- [36] Gadella M., Lara L.P., Pronko G.P. Iterative solution of some nonlinear differential equations. *Appl. Math. Comp.* 2011; 217 (22): 9480-9487.