



L'INTELLIGENCE ARTIFICIELLE AU SERVICE DU LASER INDUSTRIEL

Nous nous sommes entretenus avec MM. Rafael Barcos de Mister-Laser et Jacques Progin de Forthwood sur leurs dernières nouveautés en matière de laser industriel appuyé par de l'intelligence artificielle.

La société Forthwood est spécialisée dans la conception et la réalisation de logiciels industriels dont Forbeam, le logiciel utilisé par les têtes de travail scanner au laser de la société Mister-Laser. Cette dernière s'occupe pour sa part des contacts avec les clients ainsi que de l'installation des systèmes. Cette fructueuse coopération a débuté il y a bientôt 8 ans et permet aux deux entreprises d'être présentes à travers le monde, avec néanmoins une concentration de leurs activités en Europe.

Comment vous est venue l'idée de travailler ensemble ?

JP : Nous avons déjà travaillé par le passé sur des produits similaires et notre collaboration avait été vraiment remarquable. Quand, quelque temps plus tard, Rafael m'a demandé si j'étais intéressé à poursuivre notre travail sur un nouveau produit répondant aux derniers standards de l'industrie, je n'ai pas hésité longtemps.

RB : Sans oublier que les compétences de Jacques en programmation sont très impressionnantes et que notre produit nécessite le top du top dans ce domaine.

Justement, vous annoncez plusieurs nouveautés dans le logiciel de votre produit. Que pouvez-vous en dire ?

JP : Pour commencer, nous avons réalisé une version du serveur d'apprentissage qui peut être installée localement chez nos clients. Le défi majeur était d'encapsuler la complexité

d'installation et la taille considérable des éléments nécessaires dans un package concis, avec une procédure d'installation des plus simples.

RB : L'une des évolutions importantes est que nos clients ne doivent plus exporter leurs images vers internet sur notre serveur d'apprentissage. Même si nous traitons toujours les informations de manière confidentielle, ils ont maintenant la garantie que leurs données ne sortent pas de leurs murs.

Un serveur d'apprentissage ? Expliquez-nous...

RB : Vous vous en souvenez certainement, Eurotec a publié l'année dernière un très intéressant article sur l'algorithme Pathabene. Cet algorithme utilise un modèle IA que le serveur d'apprentissage aura mis au point pour répondre au mieux selon le cas d'utilisations. La génération de ces modèles nécessite une grande puissance de calcul et nous avons mis en place une infrastructure centralisée afin d'éviter à nos clients de devoir acquérir un matériel onéreux.

Les choses ont changé depuis l'année dernière ?

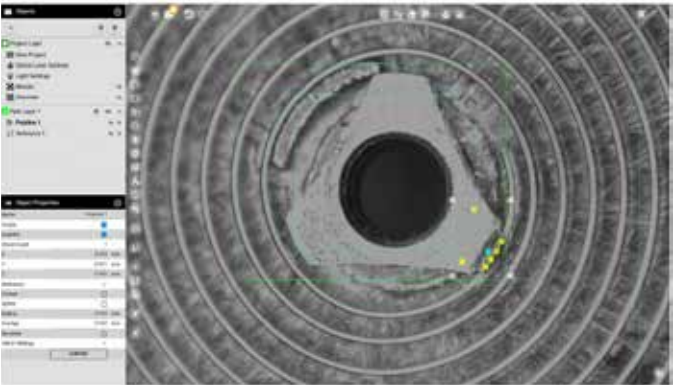
RB: Avec les dernières améliorations apportées par Jacques, ce serveur peut désormais facilement être installé directement chez nos clients. Nous avons même pu l'installer sur mon Laptop de démonstration !

Impressionnant, mais il faut certainement des spécifications très hautes pour qu'un ordinateur portable puisse recevoir le logiciel ?

JP : Pas absolument. Certes, plus elles sont hautes, plus les temps de calcul seront courts. Et comme le logiciel s'appuie sur la dernière version de PyTorch pour affiner un modèle Resnet18, il est impératif d'avoir une bonne carte graphique Nvidia avec les bibliothèques CUDA.

Abordons concrètement les performances...

RB : Nous avons installé et testé le système sous Windows 11 Pro et sous Ubuntu 24.04 LTS. L'ordinateur utilisé fut un Laptop HP Victus, Processeur 13th Gen Intel(R) Core(TM) i7-13700H 2.40 GHz, 16.0 Go de RAM et carte graphique NVidia GeForce RTX 4060 Laptop GPU. L'exercice test a consisté en l'apprentissage sur la base de 16 images de 15 Mégapixels. Cette taille d'image est parmi les plus grandes que nous ayons vu chez nos clients dans des applications de soudage. Pour donner un ordre de grandeur, voici à quoi cela ressemble :

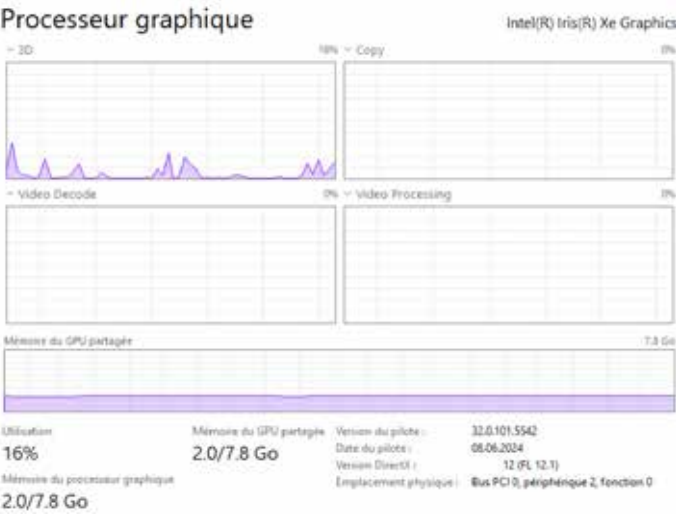


JP: En premier, nous avons fait l'apprentissage sous Windows 11 Pro et voici le temps de de travail:

State	Settings	Duration (s)	Score (%)
DONE	e: 25 lr: 0.0002 512x512	6805	78.65

Windows 11 Pro

Il ne faut pas oublier que ce temps est unique, et comme le travail est effectué sur un serveur, cela ne bloque pas la production. Le serveur travaille à créer un modèle pour Forbeam qui lui, travaille avec des temps de recherche de quelques millisecondes. Il nous a semblé étrange de constater les limitations concernant l'usage du GPU sous Windows 11 :



la capacité de calcul n'était que faiblement utilisée. Donc, en toute logique, nous avons testé le serveur sous Ubuntu 24.04 LTS sur le même ordinateur. Et les performances de calcul ont montré des GPUs utilisés à plus de 90 % et des temps radicalement inférieurs pour le même apprentissage :

State	Settings	Duration (s)	Score (%)	State	Settings	Duration (s)	Score (%)
DONE	e: 25 lr: 0.0002 512x512	1033	76.51	DONE	e: 25 lr: 0.0002 512x512	6805	78.65
Ubuntu 24.04 LTS				Windows 11 Pro			

JP : Il n'y a pas que le serveur d'apprentissage qui fonctionne sous Linux ! Nous avons également porté Forbeam pour qu'il tourne de manière native sous Linux, ce qui n'a pas été une mince affaire au vu des quelque 420'000 lignes de codes source du produit.

RB : Nous avons d'ailleurs eu l'agréable surprise de voir que Forbeam travaille entre 1.3 et 4x plus vite sous Ubuntu 24.04. LTS que sous Windows 11 Pro. Lors de notre test, le temps de calcul de l'algorithme IA Pathabene était de 1'617ms sous Windows 11 Pro et de 421ms Sous Ubuntu 24.04 LTS.

Nous avons également pu réaliser les tests sur une machine de démo dans nos locaux avec l'usage du module externe de contrôle du scanner. Ceci permet d'avoir un ordinateur portable comme unique HMI et de fortement simplifier certaines architectures de machines de micro-soudage de bureau :



C'est effectivement très compact. Il est fréquent de voir des machines imposantes dans l'industrie : pour quelles applications sont utilisés des machines aussi petites ?

RB : Notre clientèle est principalement horlogère. Le fait de pouvoir bouger manuellement l'axe Z et d'avoir un ordinateur

portable comme interface avec l'utilisateur ne requiert quasiment aucune formation. Tout a été fait pour avoir une intuitivité complète en conservant une taille de machine qui permet de la placer sur un bureau. Une prise conventionnelle de 230V suffit et aucun refroidissement externe n'est requis. Par contre, la précision de soudage est de quelques microns, même en mode manuel, comme on peut le voir sur les exemples suivants :



Merci beaucoup pour ces explications. Tout cela sera montré à l'EPHJ 2025?

RB : La plus grande partie sera effectivement à découvrir sur nos stands à l'EPHJ. Mais nous réservons également une grande surprise sur une tête de travail qui permet d'accroître le nombre de pièces pouvant être soudées au laser dans l'horlogerie, la bijouterie, l'électronique et la microtechnique par une toute nouvelle technologie laser.

A suivre....



incabloc®

décolletage - taillage - roulage



www.incabloc.ch

KÜNSTLICHE INTELLIGENZ LEISTET DEM INDUSTRIELASERBEREICH GUTE DIENSTE

Wir führten ein Gespräch mit Herrn Rafael Barcos der Firma Mister-Laser und Herrn Jacques Progin des Unternehmens Forthwood, bei dem es um KI-gestützten Industrielaser ging.

Das Unternehmen hat sich auf die Entwicklung und Herstellung von Industrie-Software-Programmen spezialisiert; auch die für die Laserscanköpfe von Mister-Laser verwendete Software Forbeam gehört dazu. Mister-Laser kümmert sich um die Kundenpflege und die Einrichtung der Systeme. Diese fruchtbare Zusammenarbeit begann vor knapp acht Jahren und ermöglicht den beiden Unternehmen, heute auf der ganzen Welt präsent zu sein, wobei sich die Geschäftstätigkeit überwiegend auf Europa konzentriert.

Wie kamen Sie auf die Idee, zusammenzuarbeiten?

JP: Wir hatten bereits in der Vergangenheit an ähnlichen Produkten gearbeitet, und unsere Zusammenarbeit lief bemerkenswert gut. Als Rafael mich einige Zeit später fragte, ob ich an der Fortsetzung unserer Zusammenarbeit interessiert sei, um ein neues Produkt, das den jüngsten Industriestandards gerecht wird, zu entwickeln, sagte ich nach kurzer Überlegung zu.

RB: Ganz abgesehen davon, dass Jacques ein hervorragender Programmierer ist und solche Kompetenzen für unser Produkt unerlässlich sind.

Sie kündigen ja gleich mehrere Neuerungen für die Software Ihres Produkts an. Können Sie uns mehr dazu sagen?

JP: Zunächst haben wir eine Lernserver-Version erstellt, die lokal bei unseren Kunden installiert werden kann. Die größte Herausforderung bestand darin, die Komplexität der Installation und den beträchtlichen Umfang der benötigten Teile geschickt zu gestalten, um ein schlankes Package mit sehr einfachen Installationsverfahren anbieten zu können.

RB: Wesentlich daran ist, dass unsere Kunden ihre Bilder nicht mehr auf unseren Lernserver im Internet exportieren müssen. Wir behandeln zwar sämtliche Informationen stets streng vertraulich, aber mit der neuen Software haben unsere Kunden die Gewähr, dass ihre Daten die eigenen vier Wände gar nicht erst verlassen.

Können Sie uns erklären, was ein Lernserver ist?

RB: Sie erinnern sich bestimmt an den Artikel über den Algorithmus Pathabene, der voriges Jahr im Eurotec-Magazin erschienen ist. Dieser Algorithmus verwendet ein vom Lernserver entwickeltes KI-Modell, um je nach Anwendungsfall bestmöglich reagieren zu können. Die Erstellung solcher Modelle erfordert eine sehr hohe Rechenleistung, und wir hatten eine zentralisierte Infrastruktur eingerichtet, um unseren Kunden die Anschaffung teurer Hardware zu ersparen.

Hat sich seit vorigem Jahr etwas geändert?

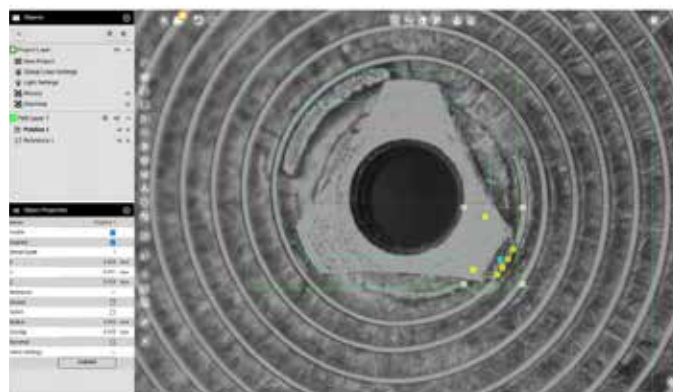
RB: Dank der von Jacques vorgenommenen Verbesserungen kann dieser Server nun problemlos direkt bei unseren Kunden installiert werden. Es ist uns sogar gelungen, ihn auf meinem Demo-Laptop zu installieren!

Das ist wirklich beeindruckend, aber für die Installation einer solchen Software ist wohl ein High-End-Laptop erforderlich ...

JP: Nicht unbedingt. Natürlich verkürzt ein schneller Computer die Rechenzeiten, und da sich die Software auf die neueste PyTorch-Version stützt, um ein Resnet18-Modell zu optimieren, ist eine gute Nvidia-Grafikkarte mit CUDA-Bibliotheken zwingend erforderlich.

Sprechen wir konkret von der Leistung ...

RB: Wir haben das System unter Windows 11 Pro und Ubuntu 24.04 LTS installiert und getestet. Wir verwendeten einen HP Victus Laptop mit einem Intel® Core™-Prozessor i7-13700H der 13. Generation, 2.40 GHz, 16.0 GB RAM und einer Grafikkarte Nvidia GeForce RTX 4060 Laptop GPU. Bei der Testübung ging es um einen Lernvorgang auf Grundlage von 16 Bildern mit 15 Megapixel. Diese Bildgröße gehört zu den größten, die wir bei unseren Kunden in Bezug auf Schweißanwendungen gesehen haben. Damit Sie sich eine Vorstellung von der Größenordnung machen können, geben wir folgendes Beispiel:



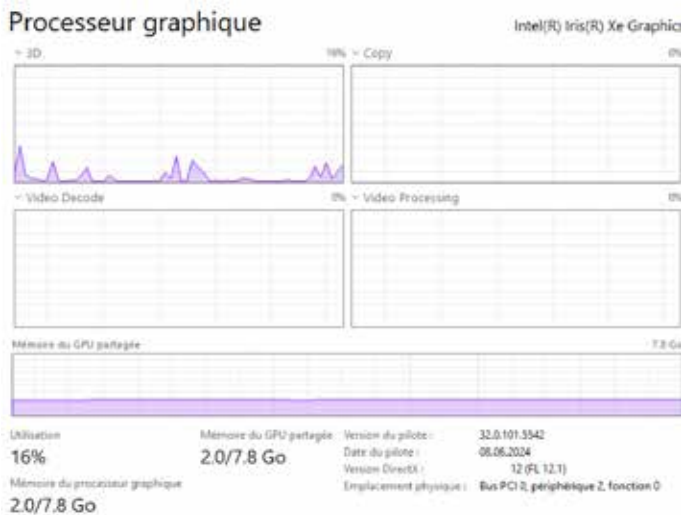
JP: Zunächst haben wir unter Windows 11 Pro gelernt, die Arbeitszeit ist nachstehend abgebildet:

State	Settings	Duration (s)	Score (%)
DONE	e: 25 lr: 0.0002 512x512	6805	78.65

Windows 11 Pro

Man darf nicht vergessen, dass diese Zeit einmalig ist, und da die Arbeit auf einem Server erledigt wird, kommt es zu keiner Behinderung der Produktion. Der Server erstellt eine Vorlage für Forbeam, der seinerseits mit einer Geschwindigkeit von ein paar Millisekunden arbeitet.

Wir wunderten uns, Einschränkungen bei der Nutzung des Grafikprozessors unter Windows 11 festzustellen:



Die Rechenkapazität wurde nur in geringem Maße genutzt. Infolgedessen haben wir den Server unter Ubuntu 24.04 LTS auf demselben Computer getestet. Die Rechenleistung zeigte GPUs, die zu mehr als 90 % ausgelastet waren, und wesentlich weniger Zeit für die gleiche Lernleistung benötigten:

State	Settings	Duration (s)	Score (%)	State	Settings	Duration (s)	Score (%)
DONE	e: 25 lr: 0.0002 512x512	1033	76.51	DONE	e: 25 lr: 0.0002 512x512	6805	78.65
Ubuntu 24.04 LTS				Windows 11 Pro			

JP: Es läuft aber nicht nur der Lernserver unter Linux! Wir haben auch Forbeam portiert, um es nativ unter Linux laufen zu lassen, was angesichts der rund 420 000 Quellcode-Zeilen des Produkts alles andere als einfach war.

RB: Wir waren sehr angenehm überrascht, als wir feststellten, dass Forbeam unter Ubuntu 24.04.LTS zwischen 1,3 und 4x schneller arbeitet als unter Windows 11 Pro. Bei unserem Test betrug die Rechenzeit des Pathabene KI-Algorithmus unter Windows 11 Pro 1'617 ms und unter Ubuntu 24.04 LTS 421 ms.

Wir konnten die Tests auch auf einer Demomaschine in unseren Räumlichkeiten durchführen, indem wir das externe Scannersteuerungsmodul verwendeten. Damit ist es möglich, einen Laptop als einzigen HMI zu verwenden und einige Desktop-Mikroschweißmaschinen-Architekturen erheblich zu vereinfachen.



Die Maschine ist in der Tat äußerst kompakt. In der Industrie ist man es gewohnt, imposante Maschinen anzutreffen: Für welche Anwendungen werden derart kleine Maschinen eingesetzt?

RB: Unsere Kunden arbeiten hauptsächlich in der Uhrenbranche. Da die Z-Achse manuell bewegt werden kann und ein Laptop als Schnittstelle zum Benutzer dient, ist so gut wie keine Schulung erforderlich. Unser Ziel war es, ein ebenso handliches wie äußerst benutzerfreundliches Gerät zu entwickeln – das ist uns auch gelungen, denn das Gerät hat auf einem Schreibtisch Platz. Eine herkömmliche 230-V-Steckdose genügt, eine externe Kühlung ist nicht erforderlich. Die Schweißgenauigkeit beträgt selbst im manuellen Modus wenige Mikrometer, wie man an den folgenden Beispielen erkennen kann:



Wir bedanken uns für diese Erklärungen. Wird das alles anlässlich der EPHJ 2025 zu sehen sein?

RB: Der Großteil wird tatsächlich auf unseren jeweiligen EPHJ-Messeständen ausgestellt sein. Allerdings behalten wir uns eine große Überraschung vor: Es geht um einen Arbeitskopf, mit dem die Anzahl der lasergeschweißten Teile in der Uhren-, Schmuck-, Elektronik- und Mikrotechnikindustrie dank einer völlig neuen Lasertechnologie erhöht werden kann.

Fortsetzung folgt ...

ARTIFICIAL INTELLIGENCE AT THE SERVICE OF INDUSTRIAL LASERS

We spoke to Rafael Barcos from Mister-Laser and Jacques Progin from Forthwood about their latest developments in industrial lasers supported by artificial intelligence.

Forthwood specialises in the design and production of industrial software, including Forbeam, the software used by Mister-Laser's laser scanner workheads. Mister-Laser is responsible for customer contacts and systems installation. This fruitful collaboration began almost 8 years ago, and now gives the two companies a worldwide presence, although their activities are concentrated in Europe.

How did you come up with the idea of working together?

JP: We'd worked together on similar products in the past, and our collaboration had been truly remarkable. When, some time later, Rafael asked me if I would be interested in continuing our work on a new product that met the latest industry standards, I didn't hesitate for long.

RB: Not to mention the fact that Jacques' programming skills are very impressive and that our product requires the best of the best in this field.

You've announced a number of new features in your product software. What can you tell us about them?

JP: To start with, we've produced a version of the learning server that can be installed locally on our customers' premises. The major challenge was to encapsulate the complexity of installation and the considerable size of the elements required in a concise package, with the simplest possible installation procedure.

RB: One important development is that our customers no longer have to export their images to the internet on our learning server. Although we still treat the information as confidential, they now have the guarantee that their data doesn't leave their walls.

A learning server? Tell us about it...

RB: As you'll remember, Eurotec published a very interesting article last year on the Pathabene algorithm. This algorithm uses an AI model that the learning server will have developed to best respond to the use case. Generating these models requires a great deal of computing power, and we had set up a centralised infrastructure to avoid our customers having to purchase expensive hardware.

Have things changed since last year?

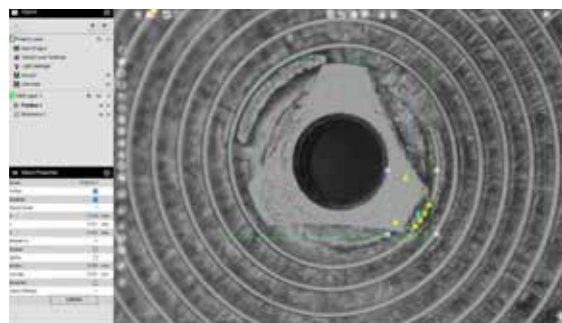
RB: With the latest improvements made by Jacques, this server can now easily be installed directly at our customers' premises. We were even able to install it on my demo Laptop!

Impressive, but surely you need a very high specification for a laptop to accommodate the software?

JP: Not absolutely. Of course, the higher they are, the shorter the calculation times will be. And as the software relies on the latest version of PyTorch to refine a Resnet18 model, it is imperative to have a good Nvidia graphics card with the CUDA libraries.

Let's take a closer look at performance...

RB: We installed and tested the system under Windows 11 Pro and Ubuntu 24.04 LTS. The computer used was an HP Victus Laptop, 13th Gen Intel(R) Core(TM) i7-13700H 2.40 GHz processor, 16.0 GB RAM and NVidia GeForce RTX 4060 Laptop GPU. The test exercise involved training on the basis of 16 15-megapixel images. This image size is one of the largest we have seen from our customers in welding applications. To give an order of magnitude, here's what it looks like:

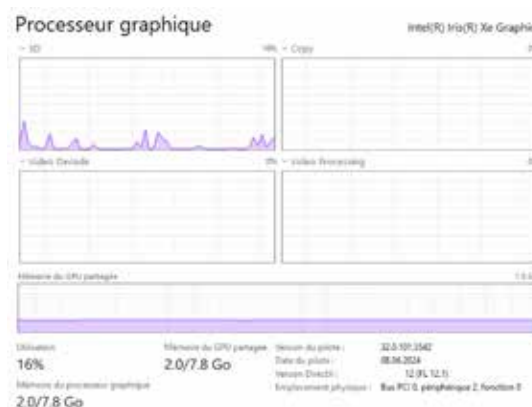


JP: First of all, we did the training on Windows 11 Pro and here's the working time:

State	Settings	Duration (s)	Score (%)
DONE	e: 25 lr: 0.0002 512x512	6805	78.65

Windows 11 Pro

It's important to remember that this time is unique, and as the work is done on a server, it doesn't block production. The server works on creating a model for Forbeam, which itself works with search times of a few milliseconds. We found it strange to note the limitations of using the GPU under Windows 11:



The computing capacity was only marginally used. So, logically, we tested the server under Ubuntu 24.04 LTS on the same

computer. And the calculation performance showed GPUs used at over 90% and radically lower times for the same training:

State	Settings	Duration (s)	Score (%)	State	Settings	Duration (s)	Score (%)
DONE	e: 25 lr: 0.0002 512x512	1033	76.51	DONE	e: 25 lr: 0.0002 512x512	6805	78.65
Ubuntu 24.04 LTS				Windows 11 Pro			

JP: It's not just the learning server that runs on Linux! We also ported Forbeam to run natively on Linux, which was no mean feat given the product's 420,000 lines of source code.

RB: We were pleasantly surprised to find that Forbeam works between 1.3 and 4x faster under Ubuntu 24.04.LTS than under Windows 11 Pro. In our test, the calculation time for the Pathabene AI algorithm was 1'617ms on Windows 11 Pro and 421ms on Ubuntu 24.04 LTS.

We were also able to carry out tests on a demo machine on our premises using the scanner's external control module. This makes it possible to use a laptop as the only HMI and greatly simplifies certain desktop micro-soldering machine architectures:



It's very compact indeed. It's common to see large machines in industry: what applications are such small machines used for?

RB: Our customers are mainly watchmakers. The fact that we can move the Z axis manually and have a laptop as the interface with the user requires virtually no training. Everything has been done to make the machine as intuitive as possible, while maintaining a size that allows it to be placed on a desk. A conventional 230V socket is all that's needed, and no external cooling is required. On the other hand, welding accuracy is within a few microns, even in manual mode, as can be seen from the following examples:



SOLUTIONS MICROTECHNIQUES SUR MESURE

**130 ans de rigueur et de précision
donnent des résultats
incomparables.**

ISO 13485:2016

P I G U E T
F R E R E S

Piguet Frères SA
Le Rocher 8
1348 Le Brassus
Switzerland

Tel. +41 (0)21 845 10 00
Fax +41 (0)21 845 10 09

info@piguet-freres.ch
www.piguet-freres.ch

- Compacte tout en offrant une précision de soudage de quelques microns.
- Kompakt und dennoch mit einer Schweißgenauigkeit von wenigen Mikronen.
- Compact, yet with a welding accuracy of just a few microns.

Thank you very much for these explanations. Will all this be on show at EPHJ 2025?

RB: Most of it will be on display on our stands at EPHJ. But we are also planning a big surprise with a workhead that will enable us to increase the number of parts that can be laser-welded in the watchmaking, jewellery, electronics and microtechnology industries using a brand new laser technology.

Stay tuned....

FORTHWOOD SÀRL
Route de Montagny 21
CH-1772 Ponthaux
T. +41 (0)26 552 02 01
www.laser.forthwood.ch

MISTER LASER
Chemin de la Chaudettaz 4
CH-1808 Les Monts-de-Corsier
T. +41 (0)21 971 11 56
www.mister-laser.com